



JOÃO ANTONIO

# **Informática para Concursos**

Teoria e Questões

5ª Edição

SÉRIE PROVAS  
& CONCURSOS

**Totalmente revisado e atualizado  
com mais de 500 imagens para  
facilitar o aprendizado.**



## DADOS DE COPYRIGHT

### Sobre a obra:

A presente obra é disponibilizada pela equipe [Le Livros](#) e seus diversos parceiros, com o objetivo de oferecer conteúdo para uso parcial em pesquisas e estudos acadêmicos, bem como o simples teste da qualidade da obra, com o fim exclusivo de compra futura.

É expressamente proibida e totalmente repudiável a venda, aluguel, ou quaisquer uso comercial do presente conteúdo

### Sobre nós:

O [Le Livros](#) e seus parceiros disponibilizam conteúdo de domínio público e propriedade intelectual de forma totalmente gratuita, por acreditar que o conhecimento e a educação devem ser acessíveis e livres a toda e qualquer pessoa. Você pode encontrar mais obras em nosso site: [LeLivros.us](#) ou em qualquer um dos sites parceiros apresentados [neste link](#)

*"Quando o mundo estiver unido na busca do conhecimento, e não mais lutando por dinheiro e poder, então nossa sociedade poderá enfim evoluir a um novo nível."*



JOÃO ANTONIO

# **Informática para Concursos**

Teoria e Questões

5ª Edição

SÉRIE PROVAS  
& CONCURSOS





Cadastre-se em [www.elsevier.com.br](http://www.elsevier.com.br) para conhecer nosso catálogo completo, ter acesso a serviços exclusivos no site e receber informações sobre nossos lançamentos e promoções.

Todos os direitos reservados e protegidos pela Lei nº 9.610, de 19/02/1998.

Nenhuma parte deste livro, sem autorização prévia por escrito da editora, poderá ser reproduzida ou transmitida sejam quais forem os meios empregados: eletrônicos, mecânicos, fotográficos, gravação ou quaisquer outros.

*Revisão:* Adriana Alves

*Editoração Eletrônica:* SBNigri Artes e Textos Ltda.

*Epub:* SBNigri Artes e Textos Ltda.

*Coordenador da Série:* Sylvio Motta

Elsevier Editora Ltda.

Conhecimento sem Fronteiras

Rua Sete de Setembro, 111 – 16º andar

20050-006 – Centro – Rio de Janeiro – RJ – Brasil

Rua Quintana, 753 – 8º andar

04569-011 – Brooklin – São Paulo – SP – Brasil

Serviço de Atendimento ao Cliente

0800-0265340

atendimento1@elsevier.com

ISBN: 978-85-352-7050-1

ISBN (Versão Eletrônica): 978-85-352-7051-8

**Nota:** Muito zelo e técnica foram empregados na edição desta obra. No entanto, podem ocorrer erros de digitação, impressão ou dúvida conceitual. Em qualquer das hipóteses, solicitamos a comunicação ao nosso Serviço de Atendimento ao Cliente, para que possamos esclarecer ou encaminhar a questão.

Nem a editora nem o autor assumem qualquer responsabilidade por eventuais danos ou perdas a pessoas ou bens, originados do uso desta publicação.

CIP-BRASIL. CATALOGAÇÃO-  
NA-FONTE  
SINDICATO NACIONAL DOS  
EDITORES DE LIVROS, RJ

Informática para concursos: [teoria e questões] / João Antonio Carvalho. – Rio de Janeiro: Elsevier, 2013.  
808 p. 17x24 cm – (Provas e concursos)

ISBN 978-85-352-7050-1

1. Informática – Problemas, questões, exercícios. 2. Serviço público – Brasil – Concursos. I. Título. II. Série

13-2204.

CDD: 004

CDU: 004

## Dedicatórias

Dedico este livro, de forma primordial, ao meu Deus amado, que sempre se faz presente em todos os momentos da minha vida. Misericordioso como ninguém mais é, Ele me deu, de presente, os demais merecedores desta dedicatória:

Minha mãe, dona Dija, guerreira forte e mestra meiga! Meu pai, João, exemplo de retidão e humanidade! As árvores mais perfeitas que Deus escolheu para darem este fruto tão falho. Meu Deus, muito obrigado, de todo o meu coração, por eles! Sem eles, este livro nunca seria sequer imaginado.

Meus filhos, Pedro e Mateus. Espelhos da eternidade do meu ser; razões da minha própria existência, sem os quais a minha vida pareceria com a morte eterna, sem luz e sem alegria. A vocês, meus anjos, meus guias, meus “nortes”, dedico mais este filhote!

Minha esposa, Ana, o amor da minha vida; minha cúmplice, minha amiga, minha censora (quando ultrapasso os limites do razoável), minha companheira, minha luz. Sem sua paciência e seu amor, eu nunca teria terminado nem a primeira edição deste!

Aos meus irmãos, Paula, Heitor e Tavo, que me entendem de um jeito só nosso (só nós entendemos como nos amamos, ninguém mais!). Aos meus irmãos “postigos” Cléber, Irziane, Juliene e Demétrius, que passaram a fazer parte desta família maluca, mas muito amada.

À minha segunda mãe, dona Ana, que me recebeu de braços e coração abertos no seio de sua família e hoje intercede por nós junto ao Pai que tanto a amou.

Deus, eu me sinto muito amado nesta família. Eu me sinto muito feliz convivendo com estes teus filhos! Deus, meu Deus, muito obrigado por eles! Ao Senhor, meu Pai, dedico este livro, mas por serdes misericordioso, sei que me permitirás dedicá-lo a eles também!

## Agradecimentos

Bom, a lista das pessoas às quais gostaria de agradecer é enorme, e, por isso, não caberia nestas páginas, mas aqui vão alguns dos que merecem ser destacados:

À minha família (toda) e aos meus amigos (que só não são família por questões biológicas), de quem não poderia esquecer nunca! Vocês têm sido realmente muito importantes para mim!

Ao professor Manoel Erhardt e a toda a família Espaço Jurídico (especialmente o Tiago, amigo querido), que demonstram um carinho e uma preocupação com todos os seus alunos e professores.

Ao mestre Sylvio Motta, sua esposa, Cida, e toda a sua família, pela acolhida carinhosa!

A Sirlene Lima, do Uniequipe, em São Paulo, minha amiga e irmã!

A Alexandre Naves, da Turma de Estudos, em Belo Horizonte, por sua amizade e carinho!

Ao professor Jorge Hélio, do curso Jorge Hélio, em Fortaleza. Obrigado pelo apoio e pelos conselhos!

Aos amigos Aquiles, Manoela, e a todos os que fazem parte do NUCE, o Núcleo de Concursos do Colégio Especial por tudo o que têm feito e simplesmente por serem como são: pessoas muito simples e de coração muito grande!

Aos mestres que, nesta longa estrada, tive oportunidade de conhecer e com quem pude dividir, por muitas vezes, as “salas dos professores” dos cursos por aí afora. Eles não são só mestres nas suas disciplinas, mas na vida, ensinando o valor de um sonho e a importância de persegui-lo.

Aos meus VERDADEIROS amigos do site [www.euvoupassar.com.br](http://www.euvoupassar.com.br), que acreditaram neste sonho e hoje fazem parte desta realidade revolucionária! A todos vocês, professores do EU VOU PASSAR, meu agradecimento sincero pelo filhote que vocês me ajudam a criar.

Aos meus alunos de todo o Brasil, meu agradecimento mais sincero. Sem dúvidas, meu trabalho é feito por vocês e para vocês. Aproveitem-no! Em especial a você, caro leitor. Sim! Você! Que confiou no meu trabalho e adquiriu este livro. Sei que ele fará jus à sua expectativa!

*João Antonio*

## O Autor

- Professor de Informática para concursos públicos há mais de 15 anos, ministrando cursos atualmente no Espaço Jurídico, em Recife; no Uniequipe, em São Paulo; na Turma de Estudos (Belo Horizonte) e em outras cidades do país.
- Fundador e coordenador do site [www.euvoupassar.com.br](http://www.euvoupassar.com.br) (visite-o! É revolucionário! Simplesmente o melhor site para concursos públicos do país!).
- Fundador e coordenador do site [www.beabyte.com.br](http://www.beabyte.com.br) (visite-o também, é uma proposta arrojada para melhorar a educação do país. E é gratuito!).
- Fundador e coordenador do site [www.redegir.com.br](http://www.redegir.com.br) (um sistema online de correção de redações).
- Fundador e coordenador do [www.conseguindo.com.br](http://www.conseguindo.com.br) (um site totalmente gratuito que o ajudará a “seguir” seus objetivos, acompanhando os concursos de sua preferência e informando a você sobre eles!).
- Autor de outras obras da Campus/Elsevier, como:  
*Noções de Informática para Concursos – 2a Edição*  
*Informática – Questões ESAF – 1a Edição*  
*Eu Vou Passar em Concursos – coautoria com Sylvio Motta – 1a Edição*
- Antes de tudo, porém, filho muito amado de Deus (mas isso não é privilégio meu! Você também é!).

Caro leitor,

Finalmente, meus amigos, chegamos à 5ª edição do livro **INFORMÁTICA PARA CONCURSOS**. Sim... Depois de muito tempo!

O capítulo de Hardware foi um dos que sofreram maiores mudanças, por causa, claro, das novas tecnologias que foram inseridas nestes últimos quatro (quase cinco) anos (desde a última edição).

Também alterei as versões do Windows (agora Windows 7), Word e Excel (agora na versão 2010) e os programas de Internet (programa de e-mail e navegador, com a novíssima versão do Internet Explorer 9).

Aproveito para explicar que as versões anteriores destes capítulos (como o Windows XP, o Word 2003 e Excel 2003) estarão disponíveis para download gratuitamente no [www.euvoupassar.com.br](http://www.euvoupassar.com.br) (visite o site, é revolucionário!) e no site da Editora Campus/Elsevier ([www.elsevier.com.br](http://www.elsevier.com.br)).

Que o Deus Pai, Todo-Poderoso o ilumine e faça você prosseguir nos seus estudos com fé e determinação!

Obrigado por confiar no meu trabalho.

*João Antonio*

# **Sumário**

## **Capa**

## **Folha de Rosto**

## **Cadastro**

## **Créditos**

## **Dedicatórias**

## **Agradecimentos**

## **O Autor**

## **Nota à 5ª Edição**

## **Capítulo 1 – Pequeno Histórico**

1.1. Como o computador chegou ao que é hoje

1.2. Computadores eletrônicos

## **Capítulo 2 – Hardware**

2.1. A parte física do computador

2.2. Os principais tipos de computadores

2.2.1. Conhecendo os microcomputadores

2.2.2. Aspecto externo de um computador

2.3. Como funciona o computador?

2.3.1. Componentes básicos do computador

2.3.1.1. CPU – Unidade Central de Processamento

2.3.1.2. Memórias

2.3.1.3. Dispositivos de E/S (Entrada/Saída)

2.3.1.4. Barramentos

2.3.2. Funcionamento básico do micro (finalmente)

2.3.2.1. Lidando com informações digitais

2.4. Os componentes do computador



#### 2.4.1. Microprocessador

##### 2.4.1.1. Marca e modelo

##### 2.4.1.2. Clocks (frequências)

##### 2.4.1.3. Memória cache

##### 2.4.1.4. Outras características

#### 2.4.2. Processadores da família Intel

##### 2.4.2.1. Considerações iniciais

##### 2.4.2.2. Principais modelos da Intel

##### 2.4.2.3. Termos do universo Intel

#### 2.4.3. Processadores da família AMD

##### 2.4.3.1. Considerações iniciais

##### 2.4.3.2. Os principais processadores da AMD

#### 2.4.4. Palavras finais sobre processadores

#### 2.4.5. Placa-mãe

##### 2.4.5.1. Controlador de memória integrado à CPU

#### 2.4.6. Memórias

##### 2.4.6.1. Classificação pelo tipo de memória

##### 2.4.6.2. Classificação das memórias por sua função no micro

#### 2.4.7. Dispositivos de entrada e saída

##### 2.4.7.1. Teclado (entrada)

##### 2.4.7.2. Monitor (saída)

##### 2.4.7.3. Mouse (entrada)

##### 2.4.7.4. Impressora (saída)

##### 2.4.7.5. Scanner (entrada)

##### 2.4.7.6. Multifuncional (entrada e saída)

##### 2.4.7.7. Modem (entrada e saída)

##### 2.4.7.8. Placa de rede (entrada e saída)

##### 2.4.7.9. Placa de rede Wi-Fi (entrada e saída)

##### 2.4.7.10. Placa de som (entrada e saída)

##### 2.4.7.11. Placa de vídeo (saída)

#### 2.4.8. Barramentos

##### 2.4.8.1. Barramento de sistema

##### 2.4.8.2. Barramentos de expansão

##### 2.4.8.3. RAID

#### 2.4.9. Fonte de alimentação

### 2.5. Considerações finais sobre hardware

## 2.6. Questões de hardware

### **Capítulo 3 – Softwares**

#### 3.1. Pequena definição sobre software

#### 3.2. Como funciona um programa?

#### 3.3. Tipos de softwares

#### 3.4. O que são arquivos?

#### 3.5. O que são pastas?

#### 3.6. Estrutura dos discos

#### 3.7. Sistema de arquivos

##### 3.7.1. FAT (Tabela de Alocação de Arquivos)

##### 3.7.2. NTFS – Sistema de arquivos do NT

#### 3.8. Processo de inicialização do computador

##### 3.8.1. Setor de boot

##### 3.8.2. Múltiplos sistemas operacionais

#### 3.9. Sistemas operacionais – conceitos

##### 3.9.1. Componentes do sistema operacional

##### 3.9.2. Tipos de sistemas operacionais

###### 3.9.2.1. Quanto à execução de programas

###### 3.9.2.2. Quanto à quantidade de usuários

##### 3.9.3. Sistemas operacionais famosos

##### 3.9.4. Linux – O “Patinho” Feio?

##### 3.9.5. iOS – Sistema operacional do iPhone e iPad

##### 3.9.6. Android – o sistema operacional da Google

### **Capítulo 4 – Microsoft Windows**

#### 4.1. Pequeno histórico do Windows

#### 4.2. Características básicas do sistema Windows

##### 4.2.1. Como o Windows entende as unidades

##### 4.2.2. Como o Windows trata os arquivos

##### 4.2.3. Extensões dos tipos de arquivos mais comuns no Windows

###### 4.2.3.1. Arquivos usados no dia a dia

###### 4.2.3.2. Arquivos de multimídia

##### 4.2.4. Atalhos

#### 4.3. Windows 7 – O mais atual

##### 4.3.1. Principais componentes do Windows 7

[4.3.1.1. Desktop \(Área de trabalho\)](#)

[4.3.1.2. Barra de tarefas](#)

[4.3.1.3. Botão Iniciar/Menu Iniciar](#)

[4.3.1.4. Barra de Tarefas \(Área dos botões e programas\)](#)

[4.3.1.5. Área de notificação \(System Tray\)](#)

[4.3.1.6. Ícones](#)

[4.3.1.7. Janelas](#)

[4.3.2. Componentes de uma janela](#)

[4.3.2.1. Barra de título](#)

[4.3.2.2. Barra de menus](#)

[4.3.2.3. Barra de ferramentas](#)

[4.3.2.4. Barra de endereço](#)

[4.3.2.5. Barra de status](#)

[4.3.3. Principais operações com janelas](#)

[4.3.3.1. Movendo uma janela](#)

[4.3.3.2. Redimensionando uma janela](#)

[4.3.3.3. Minimizando uma janela](#)

[4.3.3.4. Maximizando uma janela](#)

[4.3.3.5. Fechando uma janela](#)

[4.3.3.6. Trabalhando com várias janelas abertas](#)

[4.4. Principais programas do Windows](#)

[4.4.1. Windows Explorer](#)

[4.4.1.1. Conhecendo a interface do Explorer](#)

[4.4.1.2. Usando o Windows Explorer](#)

[4.4.1.3. Bibliotecas](#)

[4.4.2. Painel de controle](#)

[4.4.3. Acessórios do Windows](#)

[4.4.3.1. Calculadora](#)

[4.4.3.2. Bloco de notas](#)

[4.4.3.3. Wordpad](#)

[4.4.3.4. Paint](#)

[4.4.3.5. Outros acessórios do Windows 7](#)

[4.4.4. Ferramentas do sistema](#)

[4.4.4.1. Desfragmentador de disco](#)

[4.4.4.2. Monitor de Recursos](#)

[4.4.4.3 – Agendador de Tarefas](#)

[4.4.4.4. Limpeza de disco](#)

[4.4.4.5. Restauração do sistema](#)

[4.5. Outras dicas sobre o Windows](#)

[4.5.1. Combinações com a tecla \(Windows\)](#)

[4.5.2. Atributos dos arquivos](#)

[4.5.3. Windows Update](#)

[4.5.4. Comando Executar](#)

[4.5.5. Comando Desligar](#)

[4.5.6. Registro do Windows \(Registry\)](#)

[4.5.7. A Estrutura de pastas do Windows 7](#)

[4.5.7.1. Program Files e Program Files \(x86\)](#)

[4.5.7.2. Users](#)

[4.5.7.3. Windows](#)

[4.5.8. Grupo Doméstico](#)

[4.5.8.1. O que é uma Rede Doméstica?](#)

[4.6. Windows 8 – O mais novo!](#)

[4.7. Questões de Windows](#)

## **Capítulo 5 – Aplicativos Diversos**

[5.1. Conceito de aplicativos](#)

[5.1.1. Tipos de aplicativos](#)

[5.2. Instalação de um programa](#)

[5.2.1. Desinstalação de um programa](#)

[5.3. Classificação quanto à Licença de Uso](#)

[5.4. WinZip – Compactador de arquivos](#)

[5.5. Adobe Reader](#)

[5.6. Suítes de programas](#)

[5.6.1. Microsoft Office](#)

[5.6.2. LibreOffice \(antigo BrOffice\)](#)

## **Capítulo 6 – Microsoft Word**

[6.1. Conhecendo o Microsoft Word](#)

[6.1.1. Interface do Word](#)

[6.1.1.1. Faixa de Opções](#)

[6.1.1.2. Régua](#)

[6.1.1.3. Barra de status](#)

#### 6.1.1.4. Barra de rolagem

### 6.1.2. Digitando no Microsoft Word

#### 6.1.2.1. Conhecendo o texto

### 6.1.3. Selecionando trechos do texto

#### 6.1.3.1. Alguns detalhes sobre trechos selecionados

#### 6.1.3.2. Efeitos de caractere (fonte) versus efeitos de parágrafo

## 6.2. Principais comandos e recursos do Word

### 6.2.1. Guia Página Inicial

#### 6.2.1.1. Grupo Área de Transferência

#### 6.2.1.2. Grupo Fonte

#### 6.2.1.3. Grupo Parágrafo

#### 6.2.1.4. Grupo Estilo

#### 6.2.1.5. Grupo Edição

### 6.2.2. Guia Inserir

#### 6.2.2.1. Grupo Páginas

#### 6.2.2.2. Grupo Tabelas

#### 6.2.2.3. Grupo Ilustrações

#### 6.2.2.4. Grupo Links

#### 6.2.2.5. Grupo Cabeçalho e Rodapé

#### 6.2.2.6. Grupo Texto

#### 6.2.2.7. Grupo Símbolos

### 6.2.3. Guia Layout da Página

#### 6.2.3.1. Grupo Temas

#### 6.2.3.2. Grupo Configurar Página

#### 6.2.3.3. Grupo Plano de Fundo da Página

#### 6.2.3.4. Grupo Parágrafo

#### 6.2.3.5. Grupo Organizar

### 6.2.4. Guia Referências

#### 6.2.4.1. Grupo Sumário

#### 6.2.4.2. Grupo Notas de Rodapé

#### 6.2.4.3. Grupo Citações e Bibliografia

#### 6.2.4.4. Grupo Legendas

#### 6.2.4.5. Grupo Índice

### 6.2.5. Guia Correspondências

#### 6.2.5.1. Grupo Criar

#### 6.2.5.2. Demais Grupos

## 6.2.6. Guia Revisão

### 6.2.6.1. Grupo Revisão do Texto

### 6.2.6.2. Grupo Idioma

### 6.2.6.3. Grupo Comentários

### 6.2.6.4. Grupo Controle e Grupo Alterações

### 6.2.6.5. Demais Grupos

## 6.2.7. Guia Exibição

### 6.2.7.1. Grupo Modos de Exibição de Documento

### 6.2.7.2. Grupo Mostrar

### 6.2.7.3. Grupo Zoom

### 6.2.7.4. Grupo Janela

### 6.2.7.5. Grupo Macros

## 6.2.8. Ferramentas e Guias Interativas

### 6.2.8.1. Ferramentas de Imagem

### 6.2.8.2. Ferramentas de Desenho

### 6.2.8.3. Ferramentas de Tabela

### 6.2.8.4. Ferramentas de Equação

### 6.2.8.5. Ferramentas de Cabeçalho e Rodapé

## 6.3. Guia Arquivo

### 6.3.1. Comandos de Arquivo

#### 6.3.1.1. Salvar

#### 6.3.1.2. Salvar Como

#### 6.3.1.3. Abrir

#### 6.3.1.4. Fechar

### 6.3.2. Páginas da guia Arquivo

#### 6.3.2.1. Informações

#### 6.3.2.2. Recente

#### 6.3.2.3. Novo

#### 6.3.2.4. Imprimir

#### 6.3.2.5. Salvar e Enviar

#### 6.3.2.6. Ajuda

### 6.3.3. Demais comandos

#### 6.3.3.1. Opções

#### 6.3.3.2. Sair

### 6.3.4. Barra de Ferramentas de Acesso Rápido

#### 6.3.4.1. Salvar

[6.3.4.2. Desfazer](#)

[6.3.4.3. Refazer/Repetir](#)

[6.3.4.4. Estilos Rápidos](#)

[6.3.4.5. Personalizar a Barra de Ferramentas de Acesso Rápido](#)

[6.4. Questões de Word](#)

## **Capítulo 7 – Microsoft Excel**

[7.1. Conhecendo o Microsoft Excel](#)

[7.2. Interface do Excel](#)

[7.2.1. Faixa de Opções](#)

[7.2.2. Barra de Status](#)

[7.2.3. Caixa de nome](#)

[7.2.4. Barra de fórmulas](#)

[7.2.5. Guias das planilhas](#)

[7.2.6. Área da Planilha](#)

[7.2.6.1. Limites da planilha do Excel 2010](#)

[7.2.6.2. Planilha versus Pasta de Trabalho](#)

[7.3. Trabalhando com o Excel](#)

[7.3.1. Selecionando uma célula](#)

[7.3.2. Selecionando várias células](#)

[7.3.3. Inserindo dados na planilha](#)

[7.3.4. Como o Excel entende os dados](#)

[7.4. Cálculos – Automatizando o Excel](#)

[7.4.1. Operando – Operador – Operando](#)

[7.4.2. Como fazer cálculos aritméticos](#)

[7.4.3. Prioridade dos Operadores](#)

[7.4.4. Referências de Células](#)

[7.4.5. Usando a Alça de Preenchimento](#)

[7.4.6. Direção e sentido do arrasto em sequências](#)

[7.4.7. A Alça de Preenchimento para fórmulas](#)

[7.4.8. Macete para fórmulas copiadas](#)

[7.4.9. Usando referências absolutas](#)

[7.4.9.1. Usando F4 para construir referências absolutas](#)

[7.4.10. Macete na hora da prova para referências absolutas](#)

[7.4.11. Usando as funções do Excel](#)

[7.4.11.1. Funções “Intransitivas”](#)

7.4.11.2. Funções “Transitivas”

7.4.11.3. Funções “Politransitivas”

7.4.12. Usando intervalos de células

7.4.12.1. Alguns “segredos” dos intervalos de células

7.4.13. Expressões mais complexas

7.4.14. Usando funções menos comuns

7.4.14.1. Funções de Contagem

7.4.14.2. Funções de Soma Condicional

7.4.14.3. Função SE

7.5. Construindo gráficos no Excel

7.6. Outros comandos e recursos do Excel

7.6.1. Guia Página Inicial

7.6.1.1. Grupo Alinhamento

7.6.1.2. Grupo Número

7.6.1.3. Grupo Estilo

7.6.1.4. Grupo Células

7.6.1.5. Grupo Edição

7.6.2. Demais guias do Excel

7.6.2.1. Guia Inserir

7.6.2.2. Guia Layout da Página

7.6.2.3. Guia Fórmulas

7.6.2.4. Guia Dados

7.6.2.5. Guia Revisão

7.6.2.6. Guia Exibição

7.7. Valores de erros (Mensagens #)

7.8. Referência circular

7.9. Lembrando e aprimorando referências

7.9.1. Estilo de referência A1

7.9.2. Estilo de referência 3D

7.10. Considerações finais

7.11. Questões de Excel

## **Capítulo 8 – Redes de Computadores**

8.1. Conceitos iniciais

8.1.1. Classificação das redes

8.1.1.1. Quanto à extensão



### 8.1.1.2. Quanto ao funcionamento

## 8.2. Sistemas de comunicação

### 8.2.1. Classificações da transmissão

#### 8.2.1.1. Quanto ao tipo de transmissão

#### 8.2.1.2. Quanto ao sentido da transmissão

#### 8.2.1.3. Quanto à sincronização da transmissão

#### 8.2.1.4. Quanto à comutação (chaveamento) utilizada na transmissão

### 8.2.2. Problemas em uma transmissão

## 8.3. Meios físicos de transmissão

### 8.3.1. Cabo de par trançado

#### 8.3.1.1. UTP – o cabo não blindado

#### 8.3.1.2. STP – o cabo blindado

#### 8.3.1.3 Os fios do cabo de par trançado

#### 8.3.1.4. Cabo direto versus cabo cruzado

### 8.3.2. Cabo coaxial

### 8.3.3. Fibra óptica

### 8.3.4. Ondas eletromagnéticas

## 8.4. Topologias de rede

### 8.4.1. Topologia em barra (barramento)

### 8.4.2. Topologia em anel

### 8.4.3. Topologia em estrela

### 8.4.4. Topologia física versus topologia lógica

#### 8.4.4.1. Topologia lógica em estrela

#### 8.4.4.2. Topologia lógica em barramento

#### 8.4.4.3. Topologia lógica em anel

## 8.5. Um pouco mais sobre comutação de pacotes

### 8.5.1. Agora, os pacotes em si!

## 8.6. Arquiteturas de rede

### 8.6.1. Ethernet (IEEE 802.3)

#### 8.6.1.1. Como funciona a arquitetura Ethernet?

#### 8.6.1.2. CSMA/CD

#### 8.6.1.3. Conclusões sobre a arquitetura Ethernet

### 8.6.2. Token Ring (IEEE 802.5)

#### 8.6.2.1. Quadro Token (permissão)

### 8.6.3. Wi-Fi (IEEE 802.11) – Redes LAN sem fio

#### 8.6.3.1. Subpadrões 802.11

### 8.6.3.2. CSMA/CA

## 8.6.4. Segurança nas redes Wi-Fi

### 8.6.4.1. WEP (Wired Equivalent Privacy)

### 8.6.4.2. WPA (Wi-Fi Protected Access)

### 8.6.4.3. WPA 2 (IEEE 802.11i)

## 8.6.5. Mais glossário Wi-Fi

### 8.6.5.1. MIMO (Multiple-Input, Multiple-Output)

### 8.6.5.2. Hotspot

### 8.6.5.3. SSID

## 8.7. Arquiteturas para Mans e Wans

### 8.7.1. ATM

### 8.7.2. Frame Relay

### 8.7.3. WiMAX (IEEE 802.16)

### 8.7.4. IEEE 802 – redes de computadores

## 8.8. Equipamentos usados nas redes

### 8.8.1. Placa de rede (ou adaptador de rede)

#### 8.8.1.1. Endereço MAC (endereço físico)

#### 8.8.1.2. Dando outra olhada na comunicação na rede

### 8.8.2. Repetidor

### 8.8.3. Hub

#### 8.8.3.1. Hub passivo

#### 8.8.3.2. Hub ativo

### 8.8.4. Ponte

### 8.8.5. Switch

#### 8.8.5.1. Tabela de endereços MAC

### 8.8.6. Ponto de acesso (Access Point)

### 8.8.7. Roteador

#### 8.8.7.1. Tabela de roteamento

### 8.8.8. Vários componentes de rede juntos

## 8.9. Modelos de camadas

### 8.9.1. Modelo de camadas ISO/OSI

#### 8.9.1.1. Camada 1 – camada física

#### 8.9.1.2. Camada 2 – camada de enlace (ou enlace de dados)

#### 8.9.1.3. Camada 3 – camada de rede

#### 8.9.1.4. Camada 4 – camada de transporte

#### 8.9.1.5. Camada 5 – camada de sessão

[8.9.1.6. Camada 6 – camada de apresentação](#)

[8.9.1.7. Camada 7 – camada de aplicação](#)

[8.9.1.8. O modelo ISO/OSI completo](#)

[8.9.1.9. Exemplificando a comunicação \(visão do modelo OSI\)](#)

[8.9.2. Modelo de camadas TCP/IP](#)

[8.9.2.1. O “copo de liquidificador universal”](#)

[8.9.2.2. As camadas do modelo TCP/IP](#)

[8.10. Protocolos da pilha TCP/IP](#)

[8.11. Protocolos de rede](#)

[8.11.1. Protocolo IP](#)

[8.11.2. Endereço IP](#)

[8.11.3. Parâmetros IP](#)

[8.11.3.1. Endereço IP do próprio micro](#)

[8.11.3.2. Endereço IP do gateway padrão](#)

[8.11.4. Máscara de sub-rede](#)

[8.11.4.1. ID da rede e ID do host](#)

[8.11.4.2. Analisando a máscara classe C](#)

[8.11.4.3. Endereço IP da rede e endereço IP de broadcast](#)

[8.11.4.4. Analisando a máscara classe B](#)

[8.11.4.5. Analisando a máscara classe A](#)

[8.11.5. Classes de endereços IP na Internet](#)

[8.11.5.1. Então, que tal um resumo rápido?](#)

[8.11.5.2. Faixas de endereços reservados a redes privadas](#)

[8.11.6. Endereços IP especiais](#)

[8.11.7. Analisando máscaras de sub-rede binárias](#)

[8.11.8. Como o meu micro recebe os parâmetros IP?](#)

[8.11.9. IPv6 – nova forma de endereçamento na Internet](#)

[8.11.9.1. Endereços IPv6 especiais](#)

[8.11.9.2. Tipos de endereços IPv6](#)

[8.11.9.3. ID de rede e ID de host no IPv6?](#)

[8.11.10. Protocolo ICMP](#)

[8.11.11. Protocolo ARP](#)

[8.11.12. Protocolo RARP](#)

[8.12. Protocolos de transporte](#)

[8.12.1. Protocolo TCP](#)

[8.12.1.1. Estabelecimento de conexão TCP \(aperto de mãos em três vias\)](#)

[8.12.2. Protocolo UDP](#)

[8.12.3. Resumo TCP versus UDP](#)

[8.12.4. Portas](#)

[8.12.5. Socket](#)

[8.13. Protocolos de aplicação](#)

[8.13.1. SMTP](#)

[8.13.2. POP](#)

[8.13.3. IMAP](#)

[8.13.4. HTTP](#)

[8.13.5. FTP](#)

[8.13.6. Telnet](#)

[8.13.7. NNTP](#)

[8.13.8. DNS](#)

[8.13.9. DHCP](#)

[8.13.10. SNMP](#)

[8.13.11. RTP e RTCP](#)

[8.14. Outros protocolos conhecidos](#)

[8.15. Considerações finais](#)

[8.16. Questões de Redes de Computadores](#)

## **Capítulo 9 – Internet/Intranet**

[9.1. O que é a Internet?](#)

[9.2. Como a Internet nasceu?](#)

[9.3. E a Internet aqui?](#)

[9.3.1. Internet 2](#)

[9.4. Conectando-se à Internet](#)

[9.4.1. Linha telefônica \(Dial-Up\)](#)

[9.4.2. ADSL](#)

[9.4.2.1. ADSL 2/ADSL 2+](#)

[9.4.3. Internet a cabo](#)

[9.4.4. Internet através de uma rede local](#)

[9.4.5. Internet através da tomada elétrica](#)

[9.4.6. Internet via satélite](#)

[9.4.7. Internet a rádio](#)

[9.4.8. Internet via rede celular](#)

[9.5. Como funciona a Internet](#)

## 9.6. Modelo Cliente/Servidor

## 9.7. Domínios – nomes amigáveis

### 9.7.1. Hierarquia dos nomes de domínio

#### 9.7.1.1. Domínio raiz (.)

#### 9.7.1.2. Domínio de 1o nível (TLD e ccTLD)

### 9.7.2. Domínios versus nomes dos servidores

### 9.7.3. Registro de domínios no Brasil

### 9.7.4. URL – endereço único dos recursos na Internet

## 9.8. Serviços da Internet

### 9.8.1. Correio eletrônico (e-mail)

#### 9.8.1.1. Funcionamento do correio eletrônico

### 9.8.2. WWW – World Wide Web

#### 9.8.2.1. Páginas estáticas versus páginas dinâmicas

#### 9.8.2.2. Cookies

#### 9.8.2.3. Webmail

#### 9.8.2.4. Redes sociais

### 9.8.3. Transferência de arquivos – FTP

### 9.8.4. VPN – Rede Privada Virtual

### 9.8.5. Intranet

### 9.8.6. Extranet

### 9.8.7. VoIP – voz sobre IP

## 9.9. Principais aplicativos para Internet

### 9.9.1. Navegadores (Browsers)

#### 9.9.1.1. Internet Explorer

#### 9.9.1.2. Mozilla Firefox

#### 9.9.1.3. Google Chrome

#### 9.9.1.4. Outros navegadores – a galera da “geral”

### 9.9.2. Programas de correio eletrônico

#### 9.9.2.1. Mozilla Thunderbird

#### 9.9.2.2. Microsoft Live Mail

#### 9.9.2.3. Microsoft Outlook

#### 9.9.2.4. Outros programas de correio

## 9.10. Considerações finais

## 9.11. Questões de Internet

## 10.1. Comentários iniciais

## 10.2. Princípios da segurança da informação

## 10.3. Ameaças aos sistemas de informação

### 10.3.1. Malware – programas maliciosos

#### 10.3.1.1. Vírus de computador

#### 10.3.1.2. Worms

#### 10.3.1.3. Cavalos de Troia (Trojan Horses)

#### 10.3.1.4. Keyloggers e Screenloggers

#### 10.3.1.5. Spyware e Adware

#### 10.3.1.6. Backdoor (“Porta dos Fundos”)

#### 10.3.1.7. Exploits

#### 10.3.1.8. Sniffers (capturadores de quadros)

#### 10.3.1.9. Port Scanners

### 10.3.2. Fraudes e golpes na Internet

#### 10.3.2.1. Phishing (ou Phishing Scam)

#### 10.3.2.2. Pharming

#### 10.3.2.3. Engenharia social

### 10.3.3. Ataques e técnicas contra sistemas de informação

#### 10.3.3.1 Ataques DoS (Denial of Service)

#### 10.3.3.2. Buffer Overflow (sobrecarga de Buffer)

#### 10.3.3.3. Ping da morte (ping of death)

#### 10.3.3.4. SYN Flooding

#### 10.3.3.5 Spoofing

#### 10.3.3.6. Ataque Smurf

#### 10.3.3.7. Man-in-The-Middle (Homem no Meio)

## 10.4. Agentes da segurança

### 10.4.1. Antivírus

### 10.4.2. Firewall

#### 10.4.2.1. Filtro de pacotes

#### 10.4.2.2. Firewall de estado

#### 10.4.2.3. Firewall de aplicação

### 10.4.3. IDS

### 10.4.4. Antispam

### 10.4.5. DMZ – zona desmilitarizada

### 10.4.6. Bastion Host

### 10.4.7. Criptografia

## 10.5. Criptografia

### 10.5.1. Entendendo a criptografia

### 10.5.2. Criptografia é somente com números?

### 10.5.3. Criptografia simétrica (ou criptografia de chave secreta)

### 10.5.4. Entendendo a chave

### 10.5.5. Força bruta

### 10.5.6. Entendendo o tamanho da chave

#### 10.5.6.1. O usuário escolhe a sua chave?

### 10.5.7. Então, criptografia simétrica é 100% segura?

### 10.5.8. Criptografia assimétrica (criptografia de chave pública)

### 10.5.9. Finalmente, uma solução segura e funcional?

### 10.5.10. Criptografias simétrica e assimétrica juntas: um exemplo simples

### 10.5.11. A criptografia garante o quê?

## 10.6. Resumo da Mensagem (Message Digest) – Hash

### 10.6.1. As famílias de algoritmos de hash

### 10.6.2. O que obtemos com o hash?

## 10.7. Assinatura Digital

### 10.7.1. Assinatura digital na prática

### 10.7.2. O que se obtém com a assinatura digital?

### 10.7.3. A assinatura digital serve sozinha?

## 10.8. Certificação Digital

### 10.8.1. Validade do certificado

### 10.8.2. Analisando um certificado (problemas que podem ocorrer)

### 10.8.3. PKI – Public Key Infrastructure – infraestrutura de chaves públicas

### 10.8.4. E a certificação digital do ponto de vista jurídico?

### 10.8.5. A ICP-Brasil

### 10.8.6. Como emitir um certificado?

## 10.9. Então, em resumo...

## 10.10. Questões – Segurança

## **Capítulo 11 – Backup**

### 11.1. Considerações iniciais

### 11.2. Noções básicas sobre backup

### 11.3. Backups em concursos

#### 11.3.1. Para que o processo de backup é usado?

#### 11.3.2. E para que o processo de backup não é usado?

11.4. Conhecendo o processo de backup

11.5. Onde os backups são feitos?

11.6. Como os backups são feitos?

11.7. Programas para backup

11.8. A bendita “marcação” dos arquivos – teoria original de backup

11.9. Atributo de arquivamento – a “marcação” do Windows

11.10. Tipos de backups

11.10.1. Backup normal (ou global)

11.10.2. Backup incremental

11.10.3. Backup diferencial

11.10.4. Backup diário

11.10.5. Backup de cópia

11.11. Resumo dos tipos de backup

11.12. Backups cíclicos – entendendo finalmente

11.12.1. Só é importante o último ciclo

11.12.2. Estratégia 1: usando apenas backups normais

11.12.3. Estratégia 2: usando backups normais + backups incrementais

11.12.4. Estratégia 3: usando backups normais + backups diferenciais

11.12.5. Estratégia 4: usando backups normais + incrementais + diários

11.12.6. Estratégia 5: usando backups normais + diferenciais + diários

11.13. Termos e trechos “especiais” nas provas

11.14. Palavras finais

11.15. Questões de Backup

## **Capítulo 12 – Aritmética Computacional**

12.1. Sistemas numéricos

12.2. Como é formado um número

12.2.1. Entendendo alguns detalhes importantes

12.3. Processo de conversão de bases

12.3.1. Da base decimal para qualquer outra base

12.3.1.1. Exemplo prático da conversão da base decimal (endereço IP)

12.3.2. De qualquer outra base para a base decimal

12.3.3. Da base binária para a octal (e vice-versa)

12.3.4. Da base binária para a hexadecimal (e vice-versa)

12.3.5. Da base octal para a hexadecimal (e vice-versa)

12.3.6. Atenção ao índice do número



12.3.7. Operações aritméticas em bases diferentes

12.4. Operações lógicas na base 2 (noções de álgebra booleana)

12.4.1. Operador AND (E)

12.4.2. Operador OR (OU)

12.4.3. Operador NOT (Não)

12.4.4. Operador XOR (ou exclusivo)

12.4.5. Algumas regras gerais

12.4.6. Aplicação prática da álgebra booleana (endereço IP)

12.5. Considerações finais

12.6. Questões de Aritmética Computacional

**Gabaritos**

**Bibliografia**

## Capítulo 1

### Pequeno Histórico

#### 1.1. Como o computador chegou ao que é hoje

Computar significa contar, calcular, obter resultados. O computador é apenas o equipamento criado pelo homem para ajudá-lo nessa tarefa muitas vezes árdua.

Vários autores discordam sobre diversas questões filosóficas da informática, mas a grande maioria considera que o ábaco (mostrado na figura a seguir) é o primeiro computador da história da humanidade.



*Figura 1.1 – O ábaco é um aparelho que ajuda a realizar cálculos matemáticos.*

Com a contínua evolução do poder criativo e técnico da humanidade, outros dispositivos foram adicionados à história dos computadores em locais e tempos diversos. Chamo a atenção para a “máquina de cálculos” (também chamada Pascaline) criada pelo matemático Blaise Pascal no



*Figura 1.2 – Pascaline, uma calculadora mecânica criada por Pascal.*

Ao longo da história, há muitos outros relatos de dispositivos que podem ser classificados com o conceito de “computadores”, como, por exemplo, um que todos conhecem e provavelmente possuem: o termômetro. Sim, esse aparelhinho simples é conhecido como um computador analógico, pois faz a contagem de valores contínuos.

Outros equipamentos similares de medição (medição é apenas uma contagem, portanto, é uma forma de “computar”) também são considerados computadores analógicos, como as balanças, os barômetros e até os famosos bafômetros.

Nosso foco de estudo está na outra extremidade dessa classificação: vamos estudar os computadores digitais. Os computadores digitais são os equipamentos eletrônicos que manipulam informações através de pulsos elétricos que, no conceito mais superficial, podem assumir apenas dois valores: 0 (zero) e 1 (um).

OK, mas qual é a diferença entre um equipamento elétrico e um equipamento eletrônico?

Aqueles equipamentos que utilizam a energia elétrica apenas para alimentação (para acionar seus motores e dar-lhes “vida”) são chamados elétricos. Ex.: ventilador, lâmpada.

Um equipamento eletrônico (pode ser uma TV, um rádio, um computador) é um dispositivo que se alimenta da mesma energia e a manipula de forma que ela nos permita obter “respostas inteligentes”. Como por exemplo, numa televisão, a eletricidade é responsável pela alimentação, mas também é “moldada” para desenhar as imagens que vemos e os sons que ouvimos, entre outras tarefas que esse equipamento executa.

## **1.2. Computadores eletrônicos**

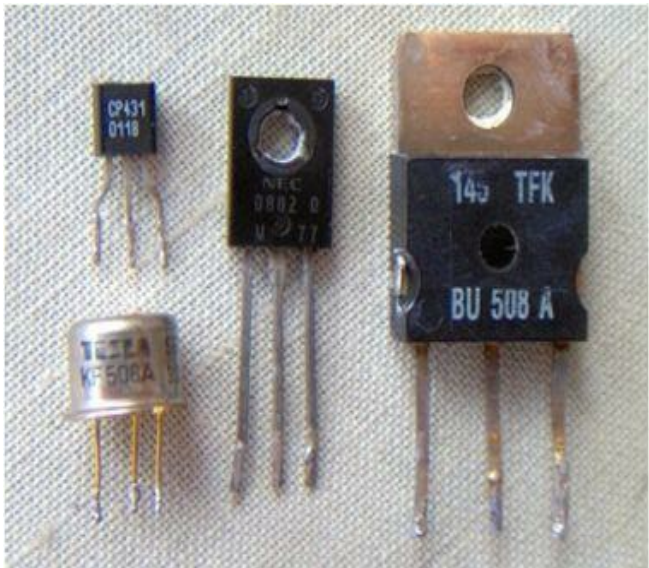
A história dos computadores eletrônicos remonta à década de 1940, quando as forças armadas dos Estados Unidos solicitaram a criação de uma máquina monstruosa, formada por milhares de válvulas (veja parágrafo seguinte) a fim de fazer cálculos importantes para a guerra. Esse computador, cujo poder de processamento é ultrapassado pelas calculadoras atuais, era chamado ENIAC.

Naquela época, manipular a energia elétrica de forma adequada era trabalho para certos componentes chamados válvulas (que, por sinal, ainda vimos aqui no Brasil em alguns modelos antigos de TVs e rádios). Portanto, todo e qualquer equipamento eletrônico (os computadores, por exemplo) tinham de ser construídos com essa tecnologia.



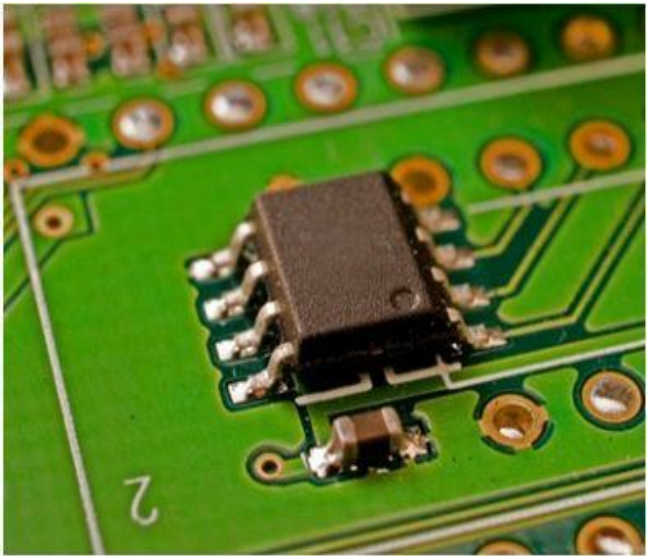
*Figura 1.3 – Válvula elétrica, uma espécie de “lâmpada” inteligente.*

Com a invenção dos semicondutores (componentes estruturais baseados no elemento químico silício), que permitem que a energia elétrica trafegue em um único sentido, foram desenvolvidas novas tecnologias, mais baratas, mais simples de usar e com resultados muito mais satisfatórios. Foi o momento de os computadores abandonarem as válvulas e usarem os transistores (Figura 1.4).



*Figura 1.4 – Transistores: componentes semicondutores usados nos computadores da segunda geração.*

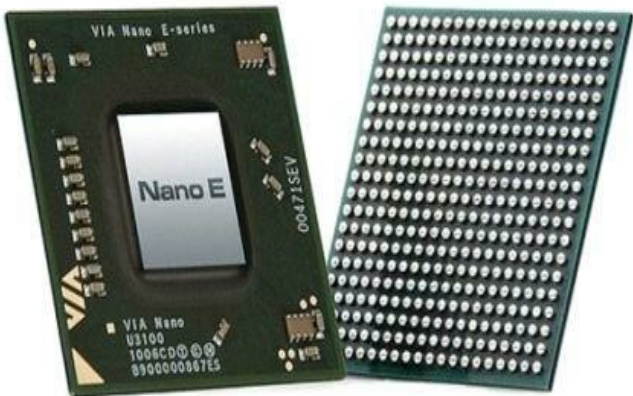
Logo após a época da grande utilização de transistores em toda a indústria de informática, desenvolveu-se uma maneira de juntar diversos desses componentes em uma única pastilha minúscula, diminuindo, em muitas vezes, o espaço necessário para montar um computador. Essas pastilhas são chamadas circuitos integrados (ou chips).



*Figura 1.5 – Chip é uma “pastilha” semicondutora que contém de centenas a milhares de transistores em seu interior. Essa é a terceira geração!*

A quarta geração foi marcada pelo uso de chips muito mais “densos” que os antecessores da terceira geração, ou seja, dentro dos chips da 4<sup>a</sup> geração havia muito mais transistores e outros componentes que nos chips da geração anterior.

Nesta geração, pode-se encontrar chips com mais de 10.000 transistores em sua composição. Essa tecnologia de miniaturização é chamada de VLSI (Escala de Integração Muito Grande). Veja na figura a seguir um exemplo de um chip da 4<sup>a</sup> Geração.



*Figura 1.6 – Chip atual, com muito mais componentes que os predecessores. São amplamente usados na indústria da informática.*

Note que não há unanimidade entre autores acerca da 5ª geração (e/ou 6ª, como alguns escrevem). O que importa é que isso não cai em concursos públicos!

Bem, esse assunto não será, precisamente, necessário para os concursos que você vai enfrentar, mas serve de base para compreender o que virá por aí! Espero que tenha gostado da Introdução... A parte boa vem agora!



## 2.1. A parte física do computador

O computador é um equipamento eletrônico utilizado para manipular informações dos mais variados tipos, como textos, fotos, desenhos, planilhas de cálculos, músicas, vídeos etc.

Nos concursos públicos atuais, há predominância de questões sobre a utilização prática do computador, ou seja, questões sobre os programas mais utilizados (programa = software). Mas isso não quer dizer que não apareçam questões sobre os componentes físicos que formam o computador, assunto estudado neste capítulo (hardware = componentes físicos, peças, dispositivos etc.).

Se o seu “xodó” é com a ESAF ou com a FGV, vá se preparando. Elas adoram o assunto de hardware! Chega a significar 20 a 30% de cada prova dessas instituições e com níveis bastante elevados. Quanto à Fundação Carlos Chagas, pode-se encontrar uma exigência maior desse assunto em provas para a área fiscal (auditores das secretarias estaduais), ou mais branda, no caso de provas para a área jurídica (técnicos e analistas de tribunais). No caso de outras instituições (Cespe/UnB, Cesgranrio, NCE/UFRJ etc.), é comum encontrar o assunto de hardware nas suas provas, mas, até agora, sem o nível de detalhamento das provas da ESAF.

## 2.2. Os principais tipos de computadores

O termo computador engloba mais que aquele equipamento que todos conhecemos, como visto a seguir:

- **Mainframe (lê-se mais ou menos assim: “meinfreîmi”):** computador de grande porte, normalmente utilizado para gerenciar grande quantidade de fluxo de dados (já imaginou quantos dados são manipulados pelos computadores centrais das operadoras de cartões de crédito? Ou dos bancos? Pois é... Aí entravam os mainframes).
- **Minicomputador:** computadores menores que os mainframes, mas ainda assim um pouco maiores que os micros que atualmente usamos. Normalmente usados em “trabalhos” que os microcomputadores não conseguiam fazer, como os trabalhos de efeitos especiais em filmes.
- **Microcomputador:** é o equipamento que todos nós conhecemos e com que estamos acostumados a lidar. Pelo fato de os concursos atualmente envolverem questões mais cotidianas, esse será nosso objetivo de estudo para este capítulo.

Note bem, amigo leitor, **hoje em dia (2013)**, os mainframes e minicomputadores deram lugar aos microcomputadores. Hoje, é possível encontrar equipamentos na escala dos micros que conseguem realizar trabalhos de altíssimo desempenho, tornando desnecessários os dispositivos das outras classificações.

“Quer dizer, João, que hoje em dia só há MICROCOMPUTADORES?”

Para a sua prova, caro leitor, **SIM!** O resto é “filosofia”...

### 2.2.1. Conhecendo os microcomputadores

O microcomputador está presente em vários momentos de nossa vida, e muitas pessoas não se veem mais sem essa utilíssima ferramenta de trabalho. Atualmente, podemos dividir os microcomputadores em vários tipos, a saber:



*Figura 2.1 – Computador desktop (um computador “de mesa”).*



*Figura 2.2 – Notebook (laptop).*

É bom registrar que **laptop** e **notebook** não são a mesma coisa. Lap é “colo” em inglês. Os laptops são um pouco maiores que os notebooks (na verdade, nem sei o porquê das duas classificações). Mas, no cotidiano, e nas provas também, esses termos hoje são tidos como sinônimos! Hoje são comuns os laptops com telas de 13, 15 e 17 polegadas.

Hoje, ainda é possível encontrar os **Ultrabooks** (notebooks muito finos e leves) que normalmente consomem menos energia (ou seja, possuem muita autonomia de bateria, permitindo que sejam usados por mais tempo sem que seja necessário recarregar a bateria do equipamento). Normalmente, usam telas de 13 polegadas.

Os ultrabooks sacrificam alguns dispositivos comuns, como drive de DVD (os ultrabooks normalmente não vêm com esse drive), para que possam ser fabricados em um chassi muito fino e elegante.

Há que se mencionar, também, os **Netbooks** (notebooks pequenos – com telas de 7, 9, 10 ou 11 polegadas), montados com processadores (cérebros) menos potentes. Essa categoria de equipamento foi desenvolvida basicamente para o acesso à Internet, contendo o básico necessário para que essa conexão aconteça.

Os netbooks são normalmente dotados de **SSD** (disco de armazenamento sólido – circuitos) em vez de **HD** (disco rígido magnético – calma, vamos ver tudo isso!) e também não vêm com

drives de DVD. Tudo isso para que o tamanho seja pequeno e que a bateria aguarde mais tempo.



*Figura 2.3 – Um Netbook de 7 polegadas (é muito difícil digitar aqui!).*

Mas, sem dúvida alguma, o computador mais “vanguardista” da atualidade é o **tablet**. Trata-se de um computador montado num chassi único, sem teclados ou mouses, com uma tela sensível ao toque.

O usuário interage diretamente na tela, onde, inclusive, aparece um “teclado” quando ele precisa digitar algo.

O maior representante deste segmento é, como não poderia deixar de ser, o **iPad** (marca registradíssima), da empresa Apple. Ele serviu de “exemplo” e “sonho de consumo” para a maioria dos outros produtos neste segmento.

Preste atenção, porém, que mais adiante, no tópico sobre dispositivos de entrada e saída (periféricos), mais precisamente na parte que fala do mouse, eu explico sobre outro equipamento que também pode receber o nome de “tablet”.

Ou seja, esse tablet que estamos vendo agora é um computador! O tablet que vamos ver mais adiante, é apenas um periférico, ou seja, “parte” de um computador!



*Figura 2.4 – Tablet (Computador de mão) – Este é o iPad da Apple.*

E, para não dizer que não falei que as flores de plástico não morrem... Opa, desculpa a viagem musical... Os **smartphones** (telefones “espertos”) são dispositivos completos, que fazem muitas coisas e ainda permitem que você ligue e atenda ligações (consulte sua operadora – é necessário pagar a conta).

Smartphones podem acessar a Internet, enviar e receber e-mails, conectar em projetores para palestras e aulas, entre outras coisas...

Digitar, aqui, porém, é quase impossível (o tamanho do teclado não deixa!). Alguns smartphones possuem teclados reduzidos, como os teclados de um telefone mesmo, em que cada tecla tem a função de 3 ou 4 letras diferentes.

Outros smartphones, porém, possuem teclado **QWERTY** (teclado normal, completo, com todas as letras nas mesmas posições de um teclado de computador). Aliás, QWERTY é o nome dado ao teclado computador por causa das seis primeiras letras da fileira superior de teclas.



*Figura 2.5 – Smartphone com teclado QWERTY.*

Apesar de bastante variados, nós vamos focar nosso estudo no funcionamento e nas características dos microcomputadores pessoais mais comuns (desktops e laptops), pois ainda são esses os mais cobrados em prova.

É claro que, eventualmente, vamos nos remeter às demais categorias apresentadas, quando for julgado necessário.

### **2.2.2. Aspecto externo de um computador**

Externamente, todo computador é basicamente igual (tomemos como exemplo, claro, o desktop). Os componentes principais de um computador de mesa podem ser vistos a seguir.



*Figura 2.6 – O computador externamente (você já conhecia, não?).*

É bom lembrar, caro leitor, que aquela “caixa metálica” que faz um barulho danado e que guarda os principais componentes funcionais do computador é chamada gabinete, e não CPU, como gostam alguns técnicos de informática (e também os leigos)!

“Ei, João, é errado falar CPU? A CPU não existe, então?”

Sim, é errado falar CPU, pois existe um componente chamado CPU, que vamos conhecer exatamente agora! Ele fica dentro do gabinete, assim como outros dispositivos.

### **2.3. Como funciona o computador?**

“Ei João! Essa é fácil! A gente digita, aparece na tela. É como uma daquelas antigas (como é mesmo o nome?) máquinas de datilografia!”

Bom, não é bem assim que a coisa funciona, não! Lembre-se: “Existe muito mais entre o teclado e o monitor do que sonha nossa vã filosofia.”

#### **2.3.1. Componentes básicos do computador**

Existem muitos componentes e processos diferentes em um computador. O simples ato de pressionar uma tecla no seu teclado ou mover o mouse até que a setinha na tela se mova é recheado de detalhes, atravessando uma série de etapas em diversos componentes. Vamos

conhecer um esquema simples que descreve os principais componentes num micro.



*Figura 2.7 – Esquema simplificado de um computador.*

Aí, você, desesperado, pergunta: “Sim, João, mas quem é quem na figura? O micro é exatamente assim? Só tem isso? E os demais termos que eu li?”

Calma, leitor.

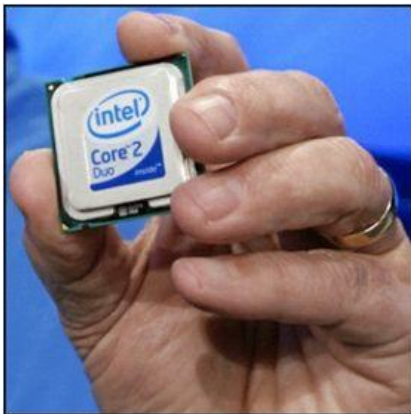
Vamos às primeiras explicações: em primeiro lugar, a figura anterior é apenas um desenho bem simplificado do computador. Com o passar do nosso assunto, vamos “aprimorando” esse esquema a fim de aproximá-lo, cada vez mais, da real “anatomia” do micro. Por ora, ficamos com isso: todo computador é composto, basicamente, por...

#### **2.3.1.1. CPU – Unidade Central de Processamento**

A CPU é simplesmente o “centro nervoso” do computador. Ela é, sem dúvida, a parte mais importante do computador. Basicamente, tudo o que se processa (processar = calcular, contar, contextualizar, transformar) em um computador é feito na CPU. Os programas que usamos, por exemplo, como o Word ou o Excel, têm suas instruções (comandos) executadas (obedecidas) pela CPU do micro.

Atualmente, as CPUs são fabricadas e comercializadas em um único componente eletrônico físico conhecido como **microprocessador**. O microprocessador (ou simplesmente **processador**) é um circuito eletrônico (chip) muito complexo e, por assumir a função de CPU, é considerado o “cérebro” do computador.





*Figura 2.8 – Um microprocessador (“micro” mesmo, hein?).*

Lembre-se de um detalhe interessante: CPU e microprocessador (ou processador) são considerados sinônimos em várias provas (e é isso o que importa para você, não é, caro leitor?). Porém, há controvérsias (e contradições) existentes em apontar a equivalência dos dois termos. (Não vamos entrar nesse mérito agora... O que importa é a prova!)

#### **2.3.1.2. Memórias**

Em primeiro lugar, é bom que se conheça o termo memória, para que se possa entender a memória principal.

Memória é, simplesmente, todo local no seu computador onde é possível **armazenar informações**. Um computador possui diversos tipos de memórias, desde as que podem guardar informações por dias, meses ou anos até aquelas que não duram muito tempo, cada qual com sua função definida.

Sim, isso mesmo! Se você pensou em CDs, DVDs e pen drives, está certo. Todos eles são memórias (não são considerados memória principal, mas são memórias). Há também o disco rígido (HD) dentro do gabinete e as memórias RAM, ROM, Cache etc. Vamos conhecê-las no momento certo!

#### **Memória Principal**

É aquela memória onde ficam guardadas as informações dos programas utilizados naquele

exato momento. A memória principal é usada não para guardar alguma coisa “para a posteridade”, mas para armazenar informações atuais: aquelas que fazem parte das janelas abertas – os programas em execução no computador.

**“Tudo o que você vê na sua tela está na memória principal”** – guarde isso!

A CPU e a MP (memória principal) se comunicam o tempo todo! Elas trocam informações constantemente enquanto o computador estiver funcionando. (Na verdade, a comunicação entre a CPU e a MP é o que faz o computador funcionar!) Se a CPU e a MP não conseguissem se comunicar, o computador nem ligaria! Vamos entender esse “caso de amor” entre elas mais adiante.

Fisicamente, a memória principal dos computadores é fabricada na forma de pequenas placas de circuitos (chamadas pentes ou módulos) contendo chips (circuitos) de um tipo de memória chamado RAM.



Figura 2.9 – Um pente (módulo) de memória RAM – usado como memória principal.

“Ei, João, espera aí! Memória principal e memória RAM são a mesma coisa, não são?”

Pergunto a você: “humano” e “aluno” são sinônimos? A resposta é NÃO! (Calma, não estou dizendo que os alunos não são humanos... Leia o resto.)

Aluno é uma função que pode ser desempenhada por indivíduos de natureza humana (*homo sapiens*) – não são sinônimos, são relacionados. Humano é natureza, existência: é “*ser*”. Aluno é função, ocupação: é “*estar*”.

**Memória principal** é uma **função**! Uma memória é chamada de principal porque é nela que são guardadas as informações utilizadas para o computador funcionar.

**Memória RAM** é um **tipo físico** de memória, uma “natureza” de memória, diferente de outras como as magnéticas (disquetes) e ópticas (CDs e DVDs). Compete à RAM assumir o papel de memória principal em nossos micros. Mais adiante trataremos nesse assunto com mais detalhamento.

### **Memórias Auxiliares**

São as memórias onde as informações conseguem ficar gravadas por tempo indeterminado (para a “posteridade”, como costumamos chamar). Essas memórias podem ter vários formatos e tamanhos.

Sim, leitor, os discos são memórias auxiliares! CDs, DVDs, HDs, disquetes e pen drives são considerados memórias auxiliares, pois mantêm as informações gravadas por muito tempo (teoricamente, até que o usuário as apague).



*Figura 2.10 – Um dispositivo de memória USB (um pen drive).*

As memórias auxiliares são também chamadas de memórias secundárias ou memórias de massa. Vamos conhecê-las com mais detalhes adiante!

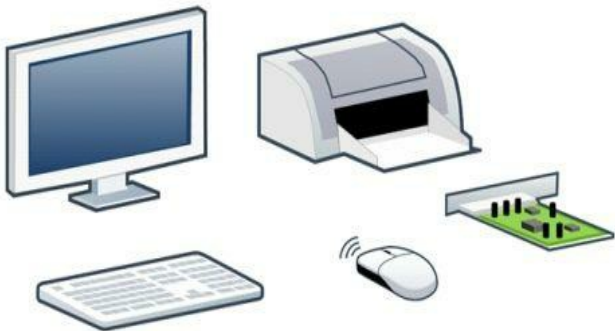
#### **2.3.1.3. Dispositivos de E/S (Entrada/Saída)**

São os equipamentos que permitem a comunicação entre a CPU e o “mundo exterior”, ou seja, o usuário. Os dispositivos de **ENTRADA** têm “mão única” e permitem a comunicação no sentido **usuário** □ **CPU**. Teclado, mouse, scanner e câmera são alguns exemplos.

Os dispositivos de **SAÍDA** também são “mão única” e permitem a comunicação no sentido **CPU** □ **usuário**. Monitor, impressora e projetor são alguns exemplos dessa classificação de equipamentos.

Há também os dispositivos **híbridos (entrada e saída)** – esses equipamentos ora permitem que informações entrem na CPU, ora permitem que elas saiam de lá. Um exemplo é o modem, responsável pela comunicação do computador com uma linha telefônica convencional, a fim de dar acesso à Internet.

Não é incomum o uso do termo “periférico” para descrever um equipamento de entrada/saída. Sim! Periféricos de entrada e periféricos de saída são expressões ainda muito utilizadas. Periférico quer dizer “aquele que está na periferia”, ou seja, “ao redor” da CPU, ajudando-a a trabalhar. Veja alguns periféricos na figura a seguir.



*Figura 2.11 – Periféricos de entrada (teclado e mouse), saída (monitor e impressora) e híbrido (modem).*

Não... Se você pensou isso, enganou-se! Os periféricos **não são** aqueles equipamentos que necessariamente estão **fora** do gabinete! Há muitos periféricos dentro do gabinete do computador, como as placas de modem (mostrada na figura anterior), som, rede e vídeo.

Novamente, como isso é apenas para “quebrar o gelo” entre você, leitor, e o assunto, os conceitos não serão muito aprofundados agora. Mais adiante veremos os periféricos com muitos

detalhes!

### 2.3.1.4. Barramentos

Note, nas figuras que apresentam o “diagrama” do micro, que há uma “estrada” interligando todos os componentes do micro. Essa via de comunicação compartilhada é chamada de barramento.

Um barramento é, em poucas palavras, um fio (ou um conjunto de fios) que funciona como uma “avenida” no micro. Sim! Há várias ruas num computador (conexões menores que não são consideradas barramentos)! Os barramentos, por sua vez, são “imponentes”, são as avenidas mais importantes.

Uma característica intrínseca ao conceito de barramento é a ideia de ***caminho compartilhado***. Os barramentos são, necessariamente, compartilhados (ou seja, não são conexões dedicadas entre dispositivos específicos).

O fato de ser compartilhado remete a um aspecto de funcionamento interessante: quando algum equipamento está utilizando o barramento, os demais não podem fazê-lo. Ou seja, apenas um equipamento pode utilizar o barramento para transmitir dados – os demais têm de esperar (ou apenas ouvir).

Você pergunta, então: “Certo, João... Mas se o barramento só pode ser usado por um dispositivo por vez, entende-se que aí há um ‘gargalo’, não é? Ou seja, se só um equipamento pode ‘falar’ por vez, o conceito de caminho compartilhado não é uma coisa tão rápida, não é? Para que serviu a ideia de barramento? E se não houvesse barramento, como seria a estrutura do micro?”

Está estudando, hein? Perguntas interessantes...

Se um computador fosse projetado sem barramentos, haveria uma série de interconexões dedicadas entre os dispositivos. Uma conexão para cada dupla de dispositivos! Seria uma conexão CPU □ memória principal; outra CPU □ discos; teclado □ CPU; CPU □ monitor e por aí vai. (Tente calcular, por um instante, o número delas.)

O barramento, por mais pontos negativos que você, leitor, encontre para ele, serve para diminuir o número de interconexões no computador, a fim de simplificar o projeto do micro. Barramentos são, em suma, recursos técnicos para tornar o micro mais simples de construir (e mais barato também).

Em um micro há vários barramentos! Cada um deles com um nome específico. Não é o momento ainda de conhecê-los, mas posso, pelo menos, mostrar-lhes as duas principais classificações de barramentos:

#### **Barramento de Sistema (System Bus)**

O barramento principal de um computador é chamado Barramento do Sistema (alguns autores o chamam de Barramento Interno). O termo System Bus também pode ser usado. O barramento de sistema interliga os principais componentes do computador (CPU, memória principal, dispositivos de E/S e memórias auxiliares).

De todos os componentes do micro, basicamente os que realmente estão ligados diretamente ao barramento de sistema são a CPU e a memória principal. Os demais componentes do micro

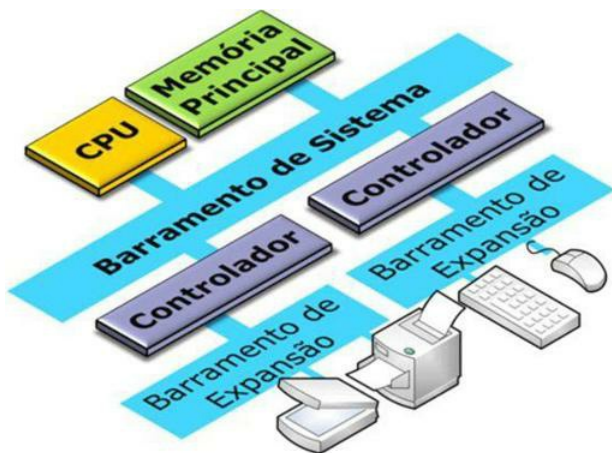
contam com dispositivos de “interface”, ou seja, dispositivos que intermedeiam a comunicação entre o barramento de sistema e os periféricos.

E antes que você reclame, a conjugação é essa mesmo: “intermedeiam”. Pode perguntar a seu professor de português!

### **Barramentos de Expansão**

São caminhos secundários, não apresentados nas imagens anteriores. Os barramentos de expansão existem para ligar periféricos ao barramento de sistema. Também podem ser chamados de barramentos de E/S (entrada e saída).

Como foi citado, os periféricos (ou dispositivos de E/S) não são ligados diretamente ao barramento de sistema, como a CPU e a MP. Eles são normalmente ligados a barramentos secundários, os quais são ligados ao barramento de sistema por meio de circuitos intermediários, normalmente conhecidos como controladores de E/S ou simplesmente controladores. A figura a seguir esclarece tudo.



*Figura 2.12 – Os barramentos de expansão e os periféricos.*

De todas as imagens que vimos até agora, a da figura anterior é a mais fiel à real estrutura do micro. O barramento de sistema e os barramentos de expansão são coisas bem distintas. Os controladores de E/S são os circuitos que gerenciam e comandam os barramentos de expansão.

Cada barramento de expansão (são vários) tem um controlador específico.

Como se você não soubesse... Mais adiante veremos detalhadamente os barramentos (tanto os de sistema quanto os de expansão).

Você deve estar pensando agora: “Ei, João... Não fuja não! Você mostrou fotos de todos os componentes até aqui! Cadê a foto dos barramentos? Onde eles estão?” – Beleza! Vamos lá!

### Onde Estão os Barramentos?

Metralhadora de perguntas para você, leitor: o barramento de sistema e os barramentos de expansão são caminhos, certo? Como se fossem “avenidas”, certo? O barramento de sistema é a avenida mais importante, como a Avenida Paulista em São Paulo, certo? Os barramentos de expansão são como “avenidas menos imponentes”, concorda?

Pois bem... E a cidade? Avenidas estão dentro de cidades, concorda? Tem de existir uma cidade. E ela existe: a **placa-mãe**. Os barramentos são parte integrante da placa-mãe.

Eu sei que a veremos detalhadamente mais adiante (novidade!), mas não custa dizer que a placa-mãe é a maior e mais importante placa de circuitos do computador. Note bem, ela não é um circuito (chip), é uma placa de circuitos (uma estrutura de resina e metal que permite a montagem de diversos chips).

Conheça a placa-mãe – a “casa” de todos os barramentos!



### **2.3.2. Funcionamento básico do micro (finalmente)**

As quatro principais etapas de funcionamento do computador são:

- 1) Entrada das informações: é o momento em que uma informação qualquer é inserida no sistema, objetivando a CPU. O exemplo mais simples é o momento da digitação de uma tecla no teclado. Claro que a entrada das informações acontece em dispositivos de entrada.
- 2) Processamento das informações: é o momento em que a CPU recebe (busca), entende (decodifica) e realiza ações (executa) com a informação que chegou. Processar é “dar destino”, transformar, contextualizar uma informação. Tá bom, sem muita “poesia”, processar é calcular! O processamento das informações, claro, é responsabilidade da CPU.
- 3) Armazenamento das informações: é a guarda dessas informações em uma memória (na maioria dos casos, incluindo o nosso exemplo atual, a memória principal). Uma vez armazenada em alguma memória, a informação já pode sair sem prejuízo ao funcionamento do micro.
- 4) Saída das informações: é o envio da informação, devidamente processada, para um dispositivo de saída (como o monitor do nosso exemplo). O simples “aparecer” na tela já é considerado uma saída de dados.

#### **2.3.2.1. Lidando com informações digitais**

Todos os dados que entram no computador, atravessam-no e são manipulados nele apresentam um formato curioso. Entendemos o que significa a letra “A” maiúscula; conseguimos compreender o sentido da expressão “34+78”, mas a máquina digital não utiliza esses códigos para realizar as operações que lhe atribuem. Ela não consegue entender tais dados da mesma maneira que nós justamente porque é digital.

Mas, o que significa ser digital? Um equipamento eletrônico pode ser analógico (como a maioria das TVs) ou digital (como os nossos computadores), mas quais as diferenças entre estas duas classificações?

Um equipamento analógico manipula a eletricidade variando-a de forma contínua... digamos, irregular. Um exemplo é quando falamos ao telefone e nossa voz é transformada em pulsos elétricos bastante irregulares que assumem diversos valores, como 0 volt, 1,2 volt, 1,3 volt, 4 volts, -3 volts, -8 volts etc. A eletricidade sendo manipulada analogicamente pode assumir qualquer valor entre o mínimo e o máximo.

Um equipamento digital faz a eletricidade variar em valores definidos, como, por exemplo, apenas entre o máximo possível e zero (um ou outro). Os nossos computadores são digitais; portanto, todas as informações manipuladas neles são consideradas apenas pulsos elétricos digitais. Verifique na figura a seguir a diferença entre esses dois tipos de informação eletrônica.



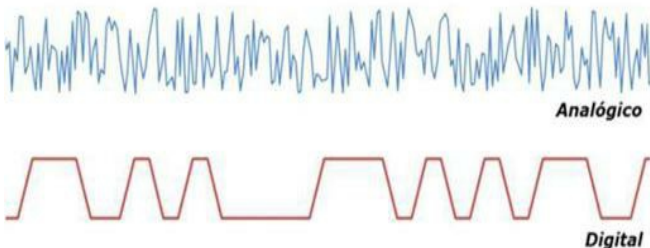


Figura 2.14 – As ondas (variações na eletricidade) digitais e analógicas são diferentes.

Como nossos computadores são máquinas digitais, vamos estudar como eles funcionam e como a **linguagem digital** (também conhecida como **binária**) está organizada.

### A Linguagem Digital – Zeros e Uns

Como vimos na figura anterior, a eletricidade em um equipamento digital só pode assumir dois valores: 0 (zero) volt, que representa a ausência de eletricidade (como uma lâmpada apagada), ou o valor máximo de voltagem do equipamento, que representa, claro, a presença de energia (como uma lâmpada acesa).

Pensando em tornar isso mais fácil de escrever e entender, foi desenvolvida a linguagem binária ou linguagem digital, que utiliza apenas dois caracteres (dois símbolos) para representar todas as informações possíveis. A linguagem binária é formada apenas por 0 (zero) e 1 (um). Que, combinados de maneiras diferentes, podem significar qualquer letra ou número que conhecemos. A seguir, um exemplo.

| Nossa Língua | Linguagem Binária |
|--------------|-------------------|
| A            | 01000001          |
| h            | 01101000          |

Para cada caractere que existe em nossos idiomas (sim, os outros países também contam!), existe um equivalente valor binário que o representa para o computador.

## **Códigos de Caracteres**

Então você pergunta: “João, e quem determina como é uma letra A maiúscula? Quem define como será representado um sinal de \$ em binário?” O engraçado, caro leitor, é que no “início dos tempos” (décadas de 1960/70), cada fabricante de computador tinha sua própria regra de definição de caracteres binários. Ou seja, cada computador tinha sua própria “língua”.

Para que não houvesse confusão de comunicação entre os diversos fabricantes e modelos de computadores, bem como as diversas línguas ao redor do mundo, foi criado um código internacional que atribui, a cada caractere nosso (ou seja, cada letra, número ou sinal), uma palavra binária (ou seja, um conjunto de zeros e uns).

Esse código, que é aceito em todo o mundo, é chamado *ASCII* (Código Americano Padrão para Intercâmbio de Informações). É ele que define, por exemplo, que a letra A é representada por 01000001.

O código ASCII usa oito zeros e uns (bits) para representar cada caractere (o que totaliza 256 combinações diferentes, já que  $2^8 = 256$ ). O Unicode, que vem sendo cotado como principal substituto do ASCII, já atribui 16 zeros e uns a cada caractere, possibilitando até 65.536 diferentes combinações em sua tabela (o que permite a representação de todos os caracteres árabes, ideogramas orientais de vários alfabetos, símbolos especiais recentes etc.). A criação do código de caracteres Unicode não aposentou o ASCII; eles convivem pacificamente nos computadores da atualidade (ou seja, seu Windows, ou Linux, nesse exato momento, é capaz de entender e-mails escritos em ASCII e em Unicode).

Lembre-se, porém, de que os códigos de caracteres que vimos são usados para representar, digitalmente, apenas texto! Outros tipos de dados de computador com que estamos acostumados, como fotos, sons e vídeos, são representados de forma binária também, mas não são baseados no ASCII nem no Unicode!

## **Representando Outros Tipos de Dados**

Só para se ter uma ideia, caro leitor, uma foto é formada por pequenos quadradinhos coloridos (os pixels, como veremos depois). Pois bem, esses quadradinhos são representados por vários zeros e uns, seguindo regras específicas que permitam ao computador interpretar aqueles conjuntos de pulsos elétricos não como letras, mas como cores!

Sim, em se tratando de texto (em ASCII), a informação 01000010 quer dizer B maiúsculo! Porém, se analisarmos a mesma informação digital como sendo parte de uma foto, esses 01000010 podem ser, digamos, um pixel verde ou um pixel amarelo... Claro! Numa música em MP3, esse mesmo 01000010 pode ser um trecho de uma nota musical tocada de um solo de guitarra, por exemplo.

Em suma, caro leitor, todas as informações que você manipula são representadas num computador na forma de 0 e 1 (digitalmente)! Dependendo, claro, do programa que se está utilizando e da informação em si, um mesmo conjunto de 0 e 1 pode ser encarado de várias maneiras!

## Bits e Bytes – Unidades de Medida

A linguagem dos computadores é digital e, portanto, baseada em 0 e 1. Seus dados são reunidos em conjuntos mínimos de oito caracteres, ou seja, oito “zeros” e “uns”. Cada caractere (cada zero ou um) é chamado *bit* (dígito binário), um conjunto de oito bits é chamado *byte* (termo binário). Veja a seguir:

**0 = bit; 1 = bit; 01100111 = byte;**

Um lembrete básico: a um conjunto de oito bits dá-se, também, o nome de *octeto*. Esse termo, porém, é mais usado em contextos como redes de computadores. A um conjunto de 4 bits (“meio byte”), dá-se o nome de *nibble* (ou semiocteto). Esse aí, sinceramente, não sei para que serve!

Para que bits e bytes são usados? No que eles podem ajudar? Esses termos representam o quê? Bytes e bits são, na verdade, unidades de medida. Assim como metro mede distância, litro mede capacidade, grama mede massa e grau Celsius mede temperatura, bytes e bits medem informação digital.

Todas as informações armazenadas em um computador são medidas em bytes. Cada texto, foto, som, desenho, filme, ou qualquer tipo de informação gravada no computador é, como vimos, um conjunto de zeros e uns, e tem seu tamanho medido não em páginas, laudas, centímetros ou minutos, mas sim, em bytes.

Um exemplo simples: a palavra **CASA** tem quatro caracteres (letras). Em ASCII, essa palavra é armazenada em **4 bytes distintos** (ou seja, ocupa 4 Bytes). Cada caractere de texto é gravado em um único byte, que pode ser visto a seguir:

**C = 01000011**

**A = 01000001**

**S = 01010011**

**A = 01000001**

Lembre-se: no ASCII, um byte é o espaço em uma memória necessário para armazenar um caractere (uma letra, por exemplo) – e essa é a ideia “geral” de armazenamento: um byte serve para armazenar uma letra. No código Unicode, para armazenar um caractere, é necessário gastar dois bytes.

Pense um pouco comigo, caro leitor: as informações digitais nos computadores são armazenadas nas memórias, não é? Portanto, todas as memórias de um computador têm sua capacidade medida em bytes. Sim: discos rígidos, DVDs, cartões de memória das máquinas fotográficas, a memória principal e até mesmo aquele seu pen drive têm suas capacidades medidas em bytes!

Lembre-se de uma coisa, leitor: a representação através de **B** (“B” maiúsculo) significa byte. O bit é representado pela letra **b** (“b” minúsculo). Algumas bancas examinadoras, porém, ainda teimam em representar, por exemplo, 30 megabytes como 30 Mb (um claro desrespeito à escrita correta dessas unidades!). A Fundação Carlos Chagas (FCC), por exemplo, já considerou, em mais de uma ocasião, siglas como mb ou Mb como sendo sinônimo de megabyte (quando, na realidade, deveria ser MB).

Lembre-se também: se alguma questão lhe fornecer um valor em bytes e requisitar o valor

correspondente em bits, basta multiplicá-lo por oito; se, por outro lado, lhe fornecerem um valor em bits e quiserem em bytes, divida-o por oito:

30 Bytes = 240 bits

32 bits = 4 Bytes

### **Kilo, Mega, Giga etc.**

Como um byte é uma unidade com valor muito pequeno, é muito comum que sejam utilizados prefixos multiplicadores conhecidos. Em medições práticas das capacidades das memórias, por exemplo, usamos a seguinte convenção de valores:

*1 Kilobyte (KB): 1.024 Bytes*

*1 Megabyte (MB): 1.024 x 1.024 Bytes*

*1 Gigabyte (GB): 1.024 x 1.024 x 1.024 Bytes*

*1 Terabyte (TB): 1.024 x 1.024 x 1.024 x 1.024 Bytes*

*1 Petabyte (PB): 1.024 x 1.024 x 1.024 x 1.024 x 1.024 Bytes*

*1 Exabyte (EB): 1.024 x 1.024 x 1.024 x 1.024 x 1.024 x 1.024 Bytes*

Dica:

*1 KB = 1.024 Bytes =  $2^{10}$  Bytes*

*1 MB = 1.024 KB =  $2^{20}$  Bytes*

*1 GB = 1.024 MB =  $2^{30}$  Bytes*

*1 TB = 1.024 GB =  $2^{40}$  Bytes*

*1 PB = 1.024 TB =  $2^{50}$  Bytes*

*1 EB = 1.024 PB =  $2^{60}$  Bytes*

Ainda há os *Zettabyte*, ou *ZB* ( $2^{70}$  Bytes) e os *Yottabyte*, ou *YB* ( $2^{80}$  Bytes).

Esses são os valores usados no nosso dia a dia pelos programas que costumamos utilizar (como o Windows e o Linux). Ou seja, esses são os valores práticos, os valores falados no cotidiano de quem usa informática na prática.

Todas essas formas de representar os múltiplos de bytes, porém, não encontram “respaldo” em documentos oficiais dos grandes órgãos padronizadores no mundo, e, por conseguinte, nas empresas que fabricam equipamentos, e é por isso que precisamos, agora, partir para uma abordagem mais “doutrinária”.

### **Kibi, Mebi, Gibi... O que é isso?**

Desde sempre, a palavra Kilo teve valor associado a 1.000 ( $10^3$ ) unidades, assim como Mega foi sempre sinônimo de 1.000.000 ( $10^6$ ), ou seja, um milhão de unidades. A maior prova disso é que se sabe que 1 km (Kilômetro) é equivalente a 1.000 metros e que 13 MW (Megawatts) é equivalente a 13.000.000 de Watts!

Todos os prefixos multiplicadores são baseados em unidades decimais, ou seja, unidades de medida escritas e calculadas numa notação matemática que tem o número 10 como base! (Nossa matemática, em suma!) Então, os prefixos Kilo, Mega, Giga, Tera etc. foram feitos para

multiplicar unidades na matemática humana! Seguem seus valores:

$$1 \text{ Kilo} = 1.000 (10^3)$$

$$1 \text{ Mega} = 1.000.000 (10^6)$$

$$1 \text{ Giga} = 1.000.000.000 (10^9)$$

$$1 \text{ Tera} = 1.000.000.000.000 (10^{12})$$

$$1 \text{ Peta} = 1.000.000.000.000.000 (10^{15})$$

$$1 \text{ Exa} = 1.000.000.000.000.000.000 (10^{18})$$

$$1 \text{ Zetta} = 1.000.000.000.000.000.000.000 (10^{21})$$

$$1 \text{ Yotta} = 1.000.000.000.000.000.000.000.000 (10^{24})$$

Portanto, considerando os verdadeiros (e tradicionais) valores das palavras Kilo, Mega, Giga etc., seria certo aceitar que 1 KB seria 1.000 bytes e não 1.024 bytes, como referido anteriormente.

O interessante é que, atualmente, **ISSO É O CERTO (para a indústria)**! Sim, pois esses são os verdadeiros valores desses prefixos no SI (Sistema Internacional de Unidades de Medida) e, consequentemente, são esses os valores praticados pelas empresas que fabricam equipamentos de memória (até mesmo porque é mais vantajoso para elas).

Note bem, apesar de eu ter mostrado os valores mais comumente aceitos para os termos Kilobyte, Megabyte e Gigabyte, bem como seus posteriores, há documentos oficiais que determinam que:

$$1 \text{ Kilobyte (KB)} = 1.000 \text{ Bytes}$$

$$1 \text{ Megabyte (MB)} = 1.000.000 \text{ Bytes (ou } 1.000 \text{ KB)}$$

$$1 \text{ Gigabyte (GB)} = 1.000.000.000 \text{ Bytes (ou } 1.000 \text{ MB)}$$

$$1 \text{ Terabyte (TB)} = 1.000.000.000.000 \text{ Bytes (ou } 1.000 \text{ GB)}$$

$$1 \text{ Petabyte (PB)} = 1.000.000.000.000.000 \text{ Bytes (ou } 1.000 \text{ TB), e assim por diante!}$$

A questão é um tanto controversa, porque bits e bytes são unidades binárias, manipuladas por equipamentos que vivem e “respiram” números binários (números na base 2). Portanto, para tais equipamentos, agrupar dígitos em conjuntos de 1.024 ( $2^{10}$ ) unidades é mais interessante (e mais simples) que reuni-los em grupos de 1.000 ( $10^3$ ) unidades.

Para que os multiplicadores usados na prática (aqueles que usam grupos de 1.024) não fossem simplesmente esquecidos, a partir da definição de 1 KB como sendo 1.000 bytes, e seus sucessores como múltiplos de 1.000, os documentos que definem o SI passaram a utilizar os seguintes prefixos para unidades binárias:

$$1 \text{ KiB (Kibibyte)} = 1.024 \text{ Bytes}$$

$$1 \text{ MiB (Mebibyte)} = 1.024 \times 1.024 \text{ Bytes (ou } 1.024 \text{ KiB)}$$

$$1 \text{ GiB (Gibibyte)} = 1.024 \times 1.024 \times 1.024 \text{ Bytes (ou } 1.024 \text{ MiB)}$$

$$1 \text{ TiB (Tebibyte)} = 1.024 \times 1.024 \times 1.024 \times 1.024 \text{ Bytes (ou } 1.024 \text{ GiB)}$$

$$1 \text{ PiB (Pebibyte)} = 1.024 \times 1.024 \times 1.024 \times 1.024 \times 1.024 \text{ Bytes (ou } 1.024 \text{ TiB);}$$

E assim vai... Existem, claro, os EiB (Exbibyte), ZiB (Zebibyte) e YiB (Yobibyte).

“Kibi? Não é aquele bolinho árabe de carne?”

Não. Kibi vem de **“Kilo-binário”**, ou seja, é o prefixo Kilo aplicado a unidades binárias (como o byte). Claro que você deduziu que Mebi significa “Mega-binário” e que Gibi é “Giga-binário”... Fácil de decorar, não?

Então, apesar de acreditarmos, pela tradição passada “de pai para filho”, que 1 Kilobyte vale 1.024 bytes, devemos ter em mente que, como há agora uma documentação oficial, usada para determinar, de uma vez por todas, os valores das unidades de medida e seus múltiplos, esses 1.024 bytes são chamados de 1 Kibibyte, pois 1 Kilobyte é o equivalente a 1.000 bytes!

### **Contagem no Windows x Contagem do Fabricante**

O fato de haver duas formas de “interpretar” os prefixos Kilo, Mega, Giga etc. tem causado uma “confusão” interessante.

Vamos imaginar uma empresa fabricante de um pen drive. Supondo que ela construa um pen drive contendo 4.000.000.000 (4 bilhões) de células de 1 Byte de memória. Logo, esse pen drive tem capacidade de 4 bilhões de Bytes, não é mesmo?

Tá, tudo bem... continuando... Se a fabricante levar em consideração que um Gigabyte é equivalente a 1.000.000.000 de bytes, então poderá, facilmente, determinar que o pen drive em questão tem **4 GB (Gigabytes)**.

Portanto, um pen drive que tem 4 bilhões de bytes de espaço é considerado, pelo fabricante, como tendo 4 GB. O problema é que o Windows (e o Linux) considera que 1 GB equivale a 1.024 x 1.024 x 1.024 bytes, e não exatamente 1 bilhão de bytes... Isso gera uma discrepância. Veja a figura a seguir:



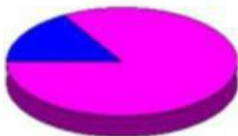
JOAO

Tipo: Disco removível

Sistema de arquivos: FAT32

|                 |                     |         |
|-----------------|---------------------|---------|
| ■ Espaço usado: | 404.922.368 bytes   | 386 MB  |
| ■ Espaço livre: | 3.594.301.440 bytes | 3,34 GB |

Capacidade: 3.999.223.808 bytes 3,72 GB



Unidade E:

Figura 2.15 – Pen drive de 4 GB... Note, o Windows diz que tem 3,72 GB.

Na figura acima, percebe-se que o pen drive tem 3.999.223.808 bytes de espaço, que significam 3,99 GB (arredondando, dá 4,0) se levarmos em consideração a interpretação de que 1 GB = 1 bilhão de bytes.

Mas o Windows lê cada GB como  $1.024 \times 1.024 \times 1.024$  bytes. Se dividirmos os 3.999.223.808 por  $1.024 \times 1.024 \times 1.024$ , obteremos 3,72! Daí o Windows afirmar que o pen drive tem 3,72 GB de capacidade (e eu garanto: quando eu comprei, na embalagem dizia 4 GB!).

Isso acontece com TODAS AS MEMÓRIAS que você compra: pen drives, discos rígidos, cartões de memória, CDs, DVDs etc. Sempre a capacidade ANUNCIADA do produto leva em consideração os múltiplos de 1.000. E o Windows (e os demais sistemas operacionais de computadores) usa a leitura de múltiplos de 1.024.

(Se você preferir assim, caro leitor, é bom entender que os sistemas operacionais Windows, Linux e outros tendem a chamar de Megabyte o que é realmente Mebibyte, e Gigabyte, o que deveria ser chamado de Gibibyte!)

Desta forma, SEMPRE haverá essa diferença entre o que é anunciado e o que efetivamente aparece na sua tela! (desista, portanto, daquela ideia de abrir uma reclamação junto ao PROCON por propaganda enganosa!)

“Ei, João... E na prova, eu digo o quê?” – Sinceramente, caro leitor, essa é uma pergunta difícil! Digo isso porque há uma dualidade aqui: o “tradicional” contra o “oficial”. E claro, até hoje, o tradicional tem vencido em provas de concursos! Isso significa que em todas as questões de provas até hoje, entendeu-se 1 kilobyte como sendo 1.024 bytes!

Cespe/UnB, FCC, ESAF, Cesgranrio... Todas elas, até hoje, demonstraram que interpretam Kilo, Mega, Giga etc. como múltiplos de **1.024**! Espero que as bancas não resolvam mudar isso justamente na sua prova!

## **2.4. Os componentes do computador**

Chegou a hora de “mergulhar” um pouco mais no estudo do hardware do computador, conhecendo, com mais detalhes, os seus principais componentes. Espero que a viagem seja agradável.

### **2.4.1. Microprocessador**

Como já foi visto anteriormente, o processador é o componente eletrônico que representa a CPU. É como se ele fosse o “cérebro” do computador. O microprocessador é, na verdade, um chip (circuito eletrônico) bastante complexo que possui a capacidade de processar (calcular) as informações que recebe. É função do microprocessador, também, executar os programas (ou seja, as instruções dos softwares são obedecidas pela CPU).

Há alguns aspectos importantes que devemos estudar relacionados a um microprocessador. Esses conceitos são amplamente exigidos em provas de concursos das mais variadas bancas examinadoras! Os aspectos mais marcantes em um processador, para análise em termos de prova de concurso, são:

- Marca e modelo;
- Clocks (frequências) interno e externo;
- Memória cache;
- Quantidade de núcleos de execução.

Ok! Ok! Concordo! Talvez o conhecimento de todas essas características não seja necessário para comprar um, mas para estudar para a prova, isso sim!

#### **2.4.1.1. Marca e modelo**

Em se tratando de processadores para computadores pessoais (os nossos), nos deparamos com duas fabricantes muito conceituadas: a Intel e a AMD. A Intel é a “casa” dos processadores Celeron, Core 2 Duo, Core i3, i5 e i7, além do Xeon. A AMD, por sua vez, é a responsável por produtos como Sempron, Phenom, Turion e Opteron.

“João, qual das duas é melhor?”



E isso lá é pergunta que se faça? Principalmente porque este livro é para concursos, que não têm histórico de confrontação entre as marcas. Seria uma pergunta realmente muito “de mau gosto” se ela exigisse a comparação técnica entre dois modelos de processadores de fabricantes diferentes.

Não vamos nos prender às marcas ainda, porque prefiro mostrar, primeiro, as características técnicas que todos os processadores apresentam para, quando analisarmos os modelos em si, possamos entender direitinho o que cada um tem de bom.

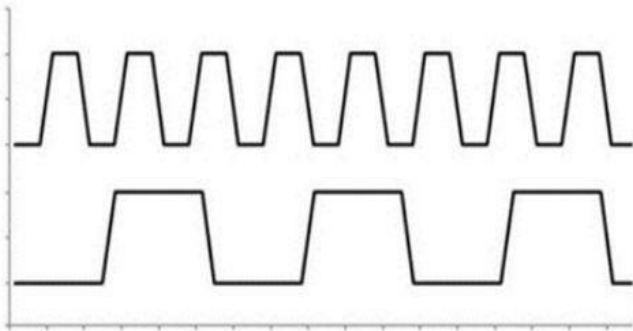
#### 2.4.1.2. Clocks (frequências)

Assim como todo equipamento eletrônico digital, os processadores possuem uma espécie de coração, que bate várias vezes por segundo. Esse “coração” é, na verdade, um pequeno cristal de quartzo que, quando alimentado de energia elétrica, gera uma onda compassada e regular, chamada **clock** (ou frequência).

É claro que, se formos muito exigentes, clock não é sinônimo de frequência. Clock (ou relógio) é o nome da onda ritmada de que estamos falando, e frequência, que, segundo o Aurélio, é o “número de ciclos que um sistema com movimento periódico efetua na unidade de tempo”, é a contagem dessa onda, a medição de suas repetições por segundo.

Os clocks dos equipamentos digitais (como o processador) têm suas frequências medidas em Hz (Hertz), a unidade internacional de frequência (lá vem esse tal de SI, de novo). 1 Hz é simplesmente 1 ciclo por segundo (ou, se você preferir, 1 acontecimento por segundo). Se algo acontece a 10 KHz, é porque se repete 10 mil vezes por segundo. Se um clock tem frequência de 800 MHz, significa que esse clock gera 800 milhões de ciclos por segundo.

Note que venho citando “frequência de clock”, ou “o clock tem tal frequência”, pois estou querendo deixar claro que são termos distintos, mas, para concursos, e para os “micreiros” de plantão, clock e frequência são o mesmo! Veja, a seguir, um exemplo de dois clocks com frequências diferentes:



*Figura 2.16 – O clock de cima tem frequência maior (repete-se mais vezes num mesmo intervalo).*

“OK, João... Muito bem... Mas para que serve um clock alto?”

Simple! Imagine que a cada ciclo (pulso completo) do clock, um determinado equipamento (no nosso caso, um processador) realiza alguma operação – o que isso te faz concluir? Quanto maior a frequência do clock (mais repetições num intervalo de tempo), mais operações aquele equipamento fará nesse intervalo.

Clocks maiores permitem a realização de mais operações por segundo. Portanto, em matéria de comparação, clocks com frequências mais altas tornam o processador mais veloz. Mas atenção para um detalhe: **clock não é velocidade!** O clock é apenas um dos fatores que determinam a velocidade final de um processador.

Não se pode, por exemplo, falar em “velocidade do processador” significando clock! É verdade que em algumas provas já se viram termos como “velocidade de clock” e “velocidade do processador”, mas eles não poderiam ser usados para descrever simplesmente a frequência.

Bem, vamos prosseguir com o nosso estudo simplesmente entendendo que todo processador tem dois clocks (duas ondas distintas com frequências distintas).

### **Clock Interno**

Dentro do processador, os componentes que o formam trabalham numa frequência bastante alta, que já ultrapassa os GHz (Gigahertz – ou bilhões de ciclos por segundo). O clock interno é usado pelos componentes internos do processador para realizar o ato do processamento das informações em si.

Hoje há processadores com 2,4 GHz, 2,8 GHz e até mesmo 3,3 GHz.

Mas uma coisa é certa! Clocks internos altos não são garantias de processadores mais rápidos.

Muitos outros fatores são importantes para se medir o “poder” de um processador. E, embora alardeado por muitos vendedores, o clock interno sequer é o aspecto mais importante para a definição da velocidade de um microprocessador.

A prova disso é que eu testei dois processadores (com várias características bem diferentes) – um com 3,2 GHz e outro com metade disso (1,6 GHz). Realizei a mesma tarefa em ambos e, para minha surpresa, o processador com 1,6 GHz de clock interno foi 6 vezes mais rápido para realizar a operação proposta. Como veem, deve haver “algo mais” que compense a frequência menor. Não se engane, leitor: ***clock interno não é tão importante assim.***

Para finalizar, se estiver sendo feita uma comparação entre dois processadores e todas as características dos dois são idênticas (com exceção do clock), é claro e evidente que aquele chip que possuir a maior frequência de clock interno será vitorioso. Ou seja, o clock interno não é o mais importante aspecto da velocidade do processador, mas como critério de desempate, claro que ele servirá.

## **Clock Externo**

O outro clock relacionado com os processadores é o clock externo. Esse é o “batimento cardíaco” que o processador usa externamente, ou seja, para se comunicar com o barramento do sistema (e, conseqüentemente, com os demais componentes do micro).

O clock externo possui uma frequência bem menor que o clock interno. Atualmente, os clocks externos variam entre 400 MHz (megahertz) a 1.066 MHz (1,06 GHz) – esses maiores, claro, presentes apenas nos processadores mais “top de linha” (leia-se: mais caros).

Antigamente, o clock externo era anunciado, em processadores da Intel como sendo ***FSB*** (Front side bus – Barramento Frontal) e pelos da AMD como sendo Barramento ***HT*** (Hyper Transport). Quando viermos a analisar os modelos dos processadores mais adiante, explicaremos isso!

## **Relação Clock Interno/Clock Externo**

Algo interessante a saber é que o clock interno e o clock externo são derivados de um clock base da placa-mãe. Sim! O processador tem, em seu interior, um cristal para gerar um clock com frequência muito alta (como vimos), mas esse cristal “não tem iniciativa” – ele simplesmente obedece à frequência do clock externo, pois, a cada pulso desta, esse cristal do processador gera ***n*** pulsos internamente.

Esse “***n***” é justamente a razão entre as duas frequências. É um valor chamado ***multiplicador***. O valor do multiplicador é definido pelo usuário, na placa-mãe, quando o micro é montado. Não é o processador que decide o multiplicador, ele apenas “obedece” às definições descritas na placa-mãe.

Lembre-se disso: o ***clock externo*** é o ritmo usado pelo processador para se comunicar com o restante do computador (em outras palavras, é o ***clock do barramento***). O clock interno, por sua vez, é o clock usado pelos componentes internos do processador (ou seja, é o clock da CPU propriamente dito), como a memória cache e os registradores, que conheceremos a seguir.

Lembre-se também de que o clock interno é derivado (múltiplo) do clock externo. Quem define o clock externo (e, por conseguinte, o interno) é o chipset da placa-mãe, que será

apresentado mais adiante.

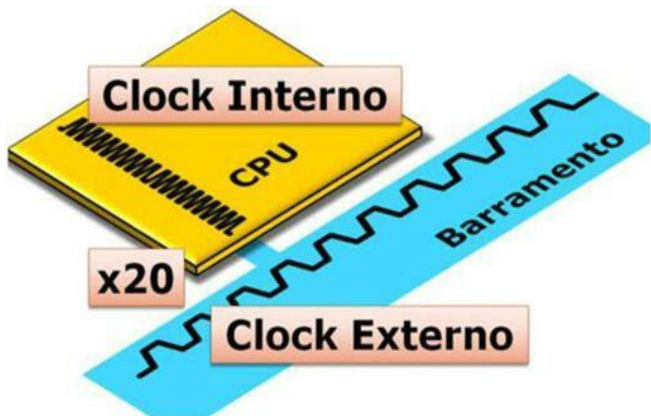


Figura 2.17 – Relação entre os clocks de um processador.

### Aumento da Frequência – Overclocking

Quando um processador é comprado, ele vem da fábrica com sua frequência interna máxima definida. Contudo, é possível alterar a frequência de seus clocks através de um processo técnico (não recomendado) chamado **overclocking**.

Esse processo consegue, com segurança, aumentos de cerca de 20% na frequência original de fábrica, em média (isso varia de modelo para modelo). Mais que isso pode fazer o processador trabalhar a uma temperatura muito superior aos limites dele, fazendo-o travar constantemente e inviabilizando seu uso.

Para realizar um overclocking, é necessário ter acesso ao programa básico que configura a placa-mãe (**setup**) e, em alguns casos, até abrir o gabinete para fazer mudanças físicas diretamente nos componentes da placa-mãe (chegando ao extremo de certos casos mais radicais em que é necessário soldar componentes à placa).

#### 2.4.1.3. Memória cache

A memória principal (como o nome já diz) é a mais importante de todas as memórias existentes no seu computador, não é? Sim, é sim... Mas não é a única.

Quando as informações são trazidas da memória principal para o processador utilizá-las, elas são depositadas, também, numa memória chamada memória cache. A cache é apenas uma memória de natureza elétrica (assim como a memória principal) que armazena informações mais rapidamente que sua amiga “importante”. A cache, na verdade, também é fabricada com chips de memória RAM (para ser mais exato, chips de memória SRAM – RAM estática –, um subtipo da RAM).

Lembre-se: cache é uma função! RAM é uma natureza física, é um tipo de memória!

As memórias cache são usadas como *memórias intermediárias* (o termo cache, em si, já está popularizado como veremos a seguir). Ela age da seguinte forma: os dados que são trazidos da memória principal são armazenados também na cache, ou seja, formando uma “cópia” deles. Quando a CPU (o processador) quiser buscar aqueles dados novamente, não será necessário consultar a memória principal (o que demora muito para a CPU): basta pedir aquela informação à cache que já a detém!

A função da cache é manter, dentro de si, o maior número possível de dados frequentemente usados, para poupar a CPU do trabalho (e gasto de tempo) excessivo em “descer” para o barramento, “baixar” de clock (para usar o externo) e ir buscar tais dados na MP. Considere a cache como, por exemplo, a lista dos 10 últimos números telefônicos para quem você ligou (comum nos celulares) – essa lista evita que o usuário precise se dirigir à agenda para achar um número telefônico para o qual sempre faz ligações.

Veja uma sequência interessante que descreve o funcionamento da cache:

É necessário, primeiramente, ter em mente que o processador e a memória principal se comunicam constantemente. Sim, em linhas gerais, a “execução” de um programa (o ato do programa funcionar) consiste numa troca constante de informações entre a MP e a CPU. Sabendo disso:

- 1) Quando a CPU precisa de uma informação, ela pergunta, primeiramente, à cache se esta possui tal informação. Se a cache já tem a informação desejada dentro de si, ela é fornecida à CPU, que a recebe e a processa. Quando a CPU encontra uma informação na memória cache, damos a isso o nome de *cache hit* (algo como “acertou na cache”).
- 2) Se a memória cache não possuía a informação solicitada pela CPU, a CPU se vê obrigada a comunicar-se com o “mundo exterior”, ou seja, com o barramento, a fim de achar a informação na memória principal. Esse é um caso de *cache miss* (que significa “faltou na cache”).
- 3) A CPU trará a informação da MP e a guardará na cache. Isso acontece para que, na próxima vez em que aquela informação for requisitada, a CPU a busque da cache, levando, para isso, muito menos tempo do que o que seria gasto se a busca fosse feita na MP.
- 4) Depois disso, a informação finalmente será entregue à CPU (considerando que houve um cache miss, claro).

Veja o resumo na figura a seguir:

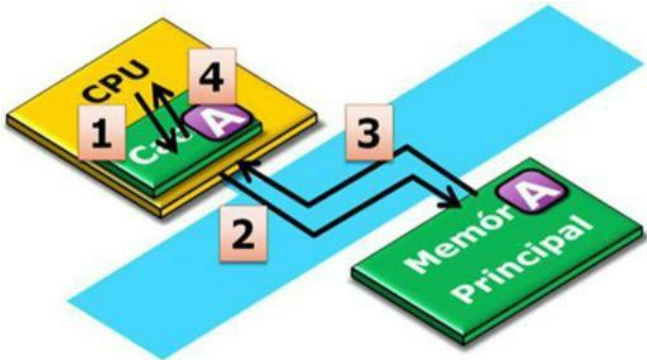
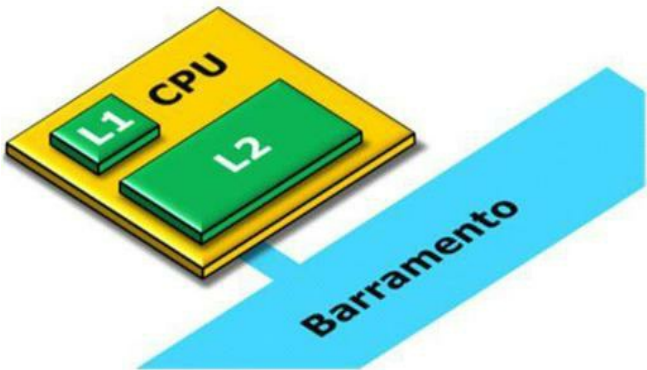


Figura 2.18 – O funcionamento da memória cache.

Note um detalhe nessa imagem:

– A memória cache faz parte do processador! Sim, atualmente, as memórias cache são fabricadas dentro do chip do processador, ou seja, na mesma pastilha de silício da CPU!

Nos processadores atuais, há duas memórias cache, ou melhor, dois níveis de memória cache: a *cache primária*, também chamada de *cache L1 (nível 1)*, e a *cache secundária*, conhecida como *cache L2 (nível 2)*. Preste muita atenção a isto: atualmente, ambos os níveis são fabricados dentro do processador!



*Figura 2.19 – Os dois níveis comuns de cache em um processador.*

“Por que a diferença? Por que dois níveis? Como funcionam?” – É fácil!

A cache L1 é a parcela da memória cache mais próxima do núcleo da CPU, construída com circuitos de memória SRAM (RAM estática) mais complexos e mais rápidos que a memória cache L2. A cache L1 é a primeira parte da cache consultada pela CPU.

Quando você ler algo a respeito de “a cache L1 tem latência mais baixa que a cache L2”, não se preocupe leitor, pois está correto! Latência diz respeito ao tempo que uma informação leva para ser transferida de um ponto a outro (no caso do nosso estudo, é o tempo que leva a informação para ir da memória cache à CPU). Como as latências das caches L1 são menores que as das caches L2, as caches L1 provam ser, sem dúvida, mais rápidas que as caches L2 (afinal, tempo menor significa velocidade maior).

E, falando em cache L2... Esse segundo nível de cache, conhecido também por cache secundária, é formado por circuitos de memória SRAM (assim como a L1) menos rápidos, até mesmo porque essa parte da cache está mais afastada (fisicamente) do núcleo da CPU.

Como os projetos dos processadores contemplam a existência de memórias cache L1 muito complexas (leia-se caras), há memórias cache L2 em maior quantidade (visto que estas são mais baratas que as L1). Enquanto as caches L1, na maioria dos processadores atuais, vão de 32 KB a 128 KB, as caches L2 já podem ser encontradas, em alguns modelos comercializados hoje, com 2 MB, 4 MB e até 8 MB (já há os “exagerados” modelos com 12 MB e até 24 MB de L2!).

Quanto mais memória cache L1 e L2 houver em um processador, melhor será o desempenho desse chip. Sim. Se houver mais espaço para que as informações mais usadas fiquem “mais próximas”, o processador levará menos tempo para achar tais informações (ou seja, será menor

a probabilidade de essa CPU precisar “se humilhar” para buscar os dados na memória principal).

Ahhh! Quase ia me esquecendo disto: em alguns processadores, existe um terceiro nível de memória cache: a cache L3 (também chamada de terciária). Esse terceiro nível de cache não é tão comum (embora esteja se tornando) nos microprocessadores dos nossos computadores pessoais, mas é bem mais provável encontrá-lo em chips de processadores mais rápidos, usados, normalmente, para computadores servidores de rede.

#### **2.4.1.4. Outras características**

##### **Conjuntos de Instruções do Processador**

Aí você pergunta: “Instruções?!”

Isso mesmo! Instruções! Um processador, mesmo sendo um equipamento físico (hardware), possui um conjunto predefinido de instruções que consegue compreender.

No conceito mais simples possível, uma instrução é uma “ordem” que o processador consegue entender, como nós, que sabemos o que significa “pegar”, “andar”, “carregar”, e outras. Todas as ordens que nos são dadas são entendidas como um conjunto sequencial de instruções simples.

Os processadores atuais (usados em nossos micros, claro!) possuem um conjunto básico de instruções conhecido como x86, que é compreendido por todos os processadores para PC. Se uma empresa qualquer fabricar um processador que não se baseie nessas instruções, ele não funcionará para PCs (ou seja, nada de Windows ou Word nele). Podemos citar como exemplos os processadores para celulares e outros dispositivos portáteis.

Com o passar do tempo, a Intel e a AMD, as principais fabricantes de processadores, desenvolveram seus “opcionais de fábrica”, vamos dizer assim. São conjuntos de instruções que só funcionam em determinados processadores e que, em alguns casos, são incompatíveis com outros.

Aqui, leitor, apenas um ALERTA! Não é necessário decorar essas datas ou detalhes dos tipos de instruções adicionadas ao longo do tempo. Basta saber seus nomes e o que elas trouxeram de novidade...

A Intel lançou o conjunto MMX (Extensões para Multimídia), para os processadores Pentium (em 1996) com a finalidade de acelerar o processamento em programas que trabalham com imagens, vídeos e som. A partir de então, todos os processadores da empresa adicionaram essas instruções ao seu conjunto básico.

Lembre-se disto: quando um conjunto de novas instruções ou recursos é adicionado a um processador, muito provavelmente todos os processadores posteriores serão dotados daquele mesmo conjunto de instruções.

A AMD lançou, para os processadores K6 e sucessores, o conjunto de instruções 3D-Now!, que continha 21 novas instruções e acelerava o processamento em programas que utilizavam ambiente tridimensional e para compactação de áudio e vídeo.

A Intel voltou a adicionar instruções aos seus modelos mais novos como o Pentium 3, que recebeu um conjunto de instruções chamado SSE para concorrer diretamente com o 3D-Now! da AMD. Os primeiros processadores Pentium 4 ganharam um conjunto de instruções chamado SSE2 que aproveita todo o poder do processador. Recentemente, para os mais novos modelos do processador Pentium 4, a Intel lançou o conjunto de instruções SSE3.



Para que as novas instruções dos processadores atuais representem um ganho significativo de desempenho, é necessário que os programas (Sistemas Operacionais, Jogos, Aplicativos etc.) sejam preparados para elas. Ou seja, não adianta comprar aquele processador CHEIO de novas instruções se os programas atuais não as utilizam.

Para que um programa esteja “apto” a utilizar perfeitamente os recursos de um processador, é necessário que este programa tenha sido programado (escrito) ou compilado (traduzido) já levando em consideração a existência de tais instruções (ou seja, o programador tem que ter usado, no meio do seu programa, algumas das instruções novas).

E lembre-se, caro leitor, novamente: não é necessário decorar todos os conjuntos de instruções que foram adicionados com o tempo. É bom entender, mas se preocupar com todos os detalhes é demais!

## **RISC x CISC**

Outra diferença de classificação entre os processadores é quanto à quantidade e o tipo de tais instruções.

As arquiteturas (ou digamos, “ideologias de funcionamento”) que dividem as opiniões dos especialistas são: CISC (Computadores com um Conjunto Complexo de Instruções) e RISC (Computadores com um Conjunto Reduzido de Instruções).

Os processadores CISC são a base da nossa computação. Os processadores que usam essa tecnologia possuem um conjunto grande de instruções, algumas delas até desnecessárias ou redundantes (repetidas), que realizam diversas operações. O problema é: quanto mais instruções um processador tem, mais complexo ele é, tornando-se, com isso, menos rápido.

As instruções presentes em processadores CISC são normalmente grandes e precisam de vários ciclos do clock do processador para serem totalmente executadas. Ou seja, quando uma ordem do tipo “escove os dentes” for dada a um processador CISC, ele a compreenderá (porque a instrução “escovar” faz parte de seu conjunto de instruções) e a executará, mas isso levará certo tempo (alguns pulsos do clock do processador).

Os processadores unicamente RISC são usados em computadores mais velozes, como algumas máquinas usadas em efeitos especiais de cinema e TV (estações de trabalho) e servidores de rede. Nesses processadores, há um conjunto de poucas instruções, apenas as mais básicas, usadas em maior quantidade. Com isso, obtém-se um ganho de desempenho considerável em relação aos processadores atuais que usamos.

Os processadores RISC conseguem, mesmo com poucas instruções, fazer o mesmo que processadores com maiores quantidades de instruções. Imagine um programador dando a um processador RISC a instrução “escove os dentes”. Como essa instrução (“escovar”) não está armazenada em seu conjunto de instruções (já que o processador RISC não tem um conjunto muito grande delas), o programa deve ser escrito (pelo programador) com instruções menores (que constam no conjunto do processador), como: “pegue a escova”, “pegue a pasta”, “ponha pasta na escova” etc.

Os processadores da tecnologia RISC executam instruções pequenas (simples), usando, na maioria dos casos, um único ciclo de clock para cada instrução.

“Ei, João, e para o usuário, qual dos dois é melhor?”

Bem, leitor, a resposta é: você é quem sabe! Afinal, é você que vai definir para que o

processador será usado e quanto se quer investir nele. Mas, uma coisa é certa: os RISC são muito mais rápidos, na maioria dos casos.

Lembre-se: os processadores de nossos computadores pessoais (Intel e AMD) são híbridos das tecnologias RISC e CISC. Ou seja, possuem um núcleo RISC, constituído de um conjunto pequeno de instruções bem rápidas, e possuem um tradutor CISC (que entende os programas construídos para CISC e quebra suas instruções para enviar ao núcleo apenas as instruções RISC simples). Eles foram construídos assim para manter a compatibilidade com os programas anteriores.

São exemplos de processadores RISC os modelos: Ultra SPARC da Sun Microsystems; Alpha da Digital e PowerPC, da IBM/Apple/Motorola (alguns deles nem mais são fabricados).

Apesar da diferença existente na forma interna como os processadores RISC e CISC trabalham, favorecendo, desde sempre, os processadores RISC, hoje em dia os nossos processadores “meio CISC, meio RISC” estão atingindo poder de processamento semelhante, ou até mesmo superior, aos processadores exclusivamente RISC (pode-se citar, por exemplo, o Core i7 e o Xeon da Intel e o Phenom e o Opteron da AMD).

### **Palavra do Processador**

Atualmente, uma das principais questões que pesam no desempenho de um processador é a sua palavra. A palavra de um processador é, basicamente, a quantidade de informação que ele pode manipular *de uma única vez*.

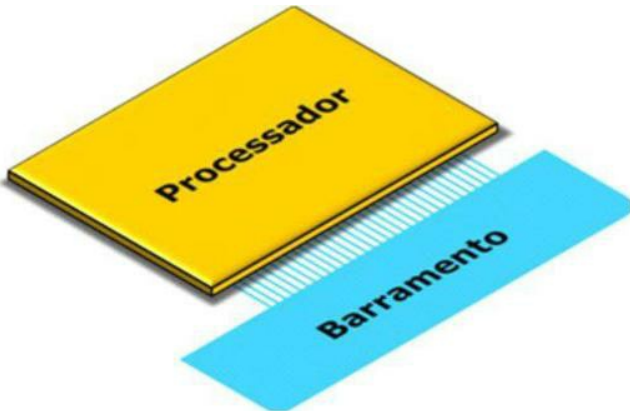
É fato que o processador é um equipamento bastante rápido, capaz de fazer inúmeras operações por segundo devido, em grande parte, ao seu clock elevadíssimo. Mas, se fosse só por causa disso, dois processadores com os mesmos clock internos e externos seriam exatamente iguais em poder de processamento, o que muitas vezes não é verdade.

Hoje em dia, a indústria da informática, especialmente no que se refere aos processadores, está focada na arquitetura de 64 bits. (Chamamos essa “medida” de PALAVRA do processador.)

“Por que a palavra é medida em bits, João?”

Essa é interessante, caro leitor! A palavra está intimamente ligada à estrutura física do computador (não só do processador, visto que a placa-mãe tem de ser compatível nesse quesito também!). A palavra é medida em bits porque está relacionada com o número de fios (que transferem 1 bit cada) que formam o barramento de dados do sistema.

Ou seja, um processador é de 32 bits quando ele pode manipular essa quantidade de informação de uma vez. E, para que isso aconteça, ou seja, para que 32 bits possam chegar de uma só vez a um processador, é claro que o barramento que transporta tais dados tem de ter, no mínimo, 32 fios (32 linhas de transmissão), como na figura a seguir.



*Figura 2.20 – Um processador ligado a um barramento de 32 bits.*

Quando um processador é de 64 bits (atualmente, todos os fabricados), o barramento de dados ligado a esse processador tem condições de transmitir 64 pulsos elétricos simultaneamente (64 bits de uma única vez). Além disso, o processador, internamente, tem condições de manipular todos esses 64 bits de uma só vez.

Mas lembre-se disto: para que um processador de 64 bits possa usar todo o seu “potencial”, é necessário que o programa que o controlará (o sistema operacional) e os programas secundários (aplicativos como o Word e o Excel) tenham sido feitos para 64 bits: é necessário que os programas tenham sido escritos (compilados) para a linguagem de máquina de 64 bits!

A principal diferença fica por conta do sistema operacional (no nosso caso, o Windows). Hoje, basicamente, todas as versões do Windows 7 (mais recente versão do programa) já são para 64 bits – ou, pelo menos, apresentam uma versão alternativa nesse formato. Em matéria de outros sistemas operacionais, como no caso do Linux, já existem versões em 64 bits há muito tempo!

Uma coisa é certa, leitor: hoje, se você for comprar uma máquina nova, excetuando os processadores mais “simples”, provavelmente já terá um computador com processador de 64 bits! Para usar tudo o que essa máquina pode dar, só com um sistema operacional feito para 64 bits (como o Windows 7 – algumas versões – e o Linux).

### **Quantidade de Núcleos de Execução**

Uma das coisas de que mais nos beneficiamos com a concorrência, além da queda nos preços dos produtos, é a busca incessante das indústrias envolvidas em desenvolver sempre o melhor

produto (ou seja, aquele que deixará a concorrente com inveja!).

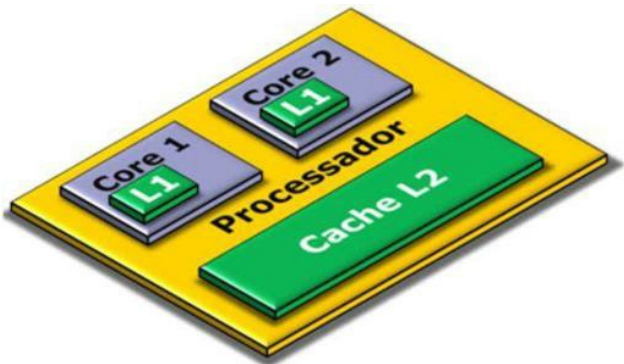
Na indústria dos cérebros de computadores, a busca por elaborar o “melhor processador” fez com que, em várias ocasiões, a Intel lançasse um produto superior ao da AMD e, em algumas semanas, recebesse a resposta à altura, acirrando ainda mais a briga.

Uma das formas que se tinha de “esquentar a briga” com a concorrente era cada vez mais aumentar o clock interno das CPUs, tornando-as mais rápidas e mais quentes (daí a razão do “esquentar”). Nesse sentido chegou-se a um limite: nenhum processador conseguiu aguentar bem, segundo relatórios das próprias empresas, clocks maiores que 4 GHz (a Intel foi a campeã, conseguindo processadores com 3,8 GHz). Algumas CPUs, durante os testes, literalmente derreteram!

Diante do limite iminente e admitindo não ter como ultrapassá-lo, buscaram-se outros artifícios para fazer CPUs mais rápidas, e, com isso, “jogar lenha na fogueira” da acirrada concorrência, sem ter, contudo, o efeito colateral da temperatura elevada: aumento da cache L1, aumento da cache L2, barramentos especiais etc. Tudo foi feito.

A melhor forma de aumentar o poder de processamento de uma CPU foi tirada, pasme leitor, de um ditado muito conhecido: “duas cabeças pensam melhor que uma”! Ou seja, a solução encontrada consiste em aumentar o número de cérebros do processador – ou número de núcleos.

Todo processador tem um conjunto central de circuitos eletrônicos microscópicos, chamado **núcleo de execução** (ou simplesmente **core** – fala-se “cór”). É nesse núcleo que são processadas (calculadas) as instruções dos programas que são executados no computador. Em outras palavras, é justamente no core que os programas (como o Windows, o Word e o Excel) realizam suas tarefas. Portanto, não é de se admirar que muitos digam que os dois núcleos são, na verdade, dois processadores (embora esse não seja o entendimento mais popular).



A briga por processadores **multicore** (muitos núcleos) começou com os processadores **dual core** (com dois núcleos), como o visto na figura anterior. Tanto a Intel quando a AMD já fabricam processadores dual core há alguns anos. Hoje, porém, os processadores que mais chamam a atenção do mercado são os que têm quatro núcleos, também chamados de **Quad core**.

Já há, também, processadores com seis núcleos (hexa core) e até oito núcleos (octa core) disponíveis no mercado, mas, como se pode deduzir, são extremamente caros.

### **Litografia (ou Tecnologia de Miniaturização) de Fabricação**

Uma das mais belicosas rixas entre as duas fabricantes de processadores está na, cada vez mais moderna, miniaturização dos componentes eletrônicos usados na fabricação dos processadores.

Claro que se sabe que um processador, hoje em dia, possui bilhões de componentes eletrônicos simples (os transistores, ou semicondutores) em sua estrutura. E fazer um produto tão complexo e denso exige um processador bem grande (em tamanho físico).

Mas o que se vê hoje é que os processadores estão diminuindo de tamanho físico e, ao mesmo tempo, aumentando de complexidade (número de componentes) – o que, em si, parece contraditório, não acha, amigo leitor?

Pois é, os componentes eletrônicos que formam os processadores também estão diminuindo de tamanho. Já saíram da faixa “microscópica” para a medição em escala de décimos de microns. Vamos às explicações:

Um micron (ou **micrômetro**) é uma unidade de medida que representa 1 metro dividido por um milhão ( $10^{-6}$  metro). Para se ter uma ideia, isso é menor que a “poeira causada pelo peido de uma pulga” (tive de usar essa medida extremamente técnica para que você pudesse perceber quão pequeno é um micron).

Pois bem, agora se fabricam componentes para processadores medindo-os em nanômetros (nm), que é nada mais que 1 metro dividido por 1 bilhão ( $10^{-9}$  metro), que seria algo em torno de “a poeira causada pelo peido de um verme dentro do intestino de uma pulga” (viu como é menor?). Quanto menor a medida, em nanômetros, dos componentes, mais moderna é a tecnologia de fabricação.

A maioria dos processadores atuais é fabricada em tecnologia de 32nm (ou, se preferir, 0,032 micron).

Sinceramente, caro leitor, acho que esse “critério” dos nanômetros não influencia muito na velocidade da máquina em si. Ou talvez influencie tudo, visto que o aumento de cache e de núcleos se deve, em parte, à redução do tamanho desses componentes. O que você acha? Dê sua opinião em [www.opeidodapulga.com](http://www.opeidodapulga.com) (brincadeira, óbvio que esse site não existe).

### **2.4.2. Processadores da família Intel**

#### **2.4.2.1. Considerações iniciais**

A Intel é a mais famosa e bem-sucedida empresa fabricante de processadores para PC do

mundo. Claro que ela não se limita a fabricar processadores, pois tem em seu rol de produtos diversos equipamentos, como placas-mãe, chipsets, placas de rede, placas de vídeo, entre outras.

Se algum concurso vier a exigir conhecimento em modelos de processadores, é muito provável que sejam mencionados modelos desta fabricante.

#### 2.4.2.2. Principais modelos da Intel

A Intel atualmente fabrica e comercializa os processadores da família “Core i”, representada pelos seus modelos Core i3, Core i5 e Core i7.

##### Intel Core i3

Criado para um público menos exigente, este processador é normalmente encontrado em ultrabooks, notebooks e desktops mais baratos.

Não conta com um poder de processamento tão alto quando seus “irmãos” mais potentes, mas é muito bom para quem irá trabalhar com operações mais simples, como digitar textos e planilhas.

O Core i3 é fabricado com dois núcleos (dual core) e possui 3 MB de Cache L2, normalmente. Há vários submodelos com clocks que variam de 2 GHz a 3.3 GHz.



*Figura 2.22 – Processador Intel Core i3.*

A segunda geração dos processadores Core i3 já é fabricada em 32nm e conta com placa gráfica (placa de vídeo) integrada ao processador (ou seja, não é necessário possuir uma placa de vídeo no computador, pois o próprio processador é capaz de gerenciar as imagens desta máquina).

Essa característica de ter a placa gráfica integrada ao processador é extremamente

interessante para notebooks (e ultrabooks, claro), pois se diminui a quantidade de equipamentos extras na placa-mãe.

### **Intel Core i5**

Produzidos para satisfazer ao mercado de micros desktop de médio porte e laptops de uso geral, os processadores Core i5 são, sem dúvidas, melhores que os processadores Core i3.

São encontrados submodelos desta família com 2 e 4 núcleos. Podem ser vistos exemplares com 3 e 6 MB de Cache L2. As faixas de clock vão de 2.1 GHz a 3.8 GHz.



*Figura 2.23 – Processador Intel Core i5.*

O processador Core i5 também possui controlador de vídeo (placa gráfica) integrado ao corpo do processador. Além disso, alguns modelos de i5 são dotados das tecnologias Intel Turbo Boost e Intel vPro (vamos falar sobre elas mais adiante).

Alguns modelos do processador Core i5 já são fabricados usando uma litografia de 22 nm (nanômetros), que é um patamar muito recente.

### **Intel Core i7 – ai, ai, Jesus...**

Eis o “sonho de consumo”. O processador Core i7 é “o carro chefe” da geração Core i. É o processador mais rápido e mais potente de todos. Foi criado para o público que realmente exige muito do computador, como quem trabalha com edição de vídeos e fotos ou quem joga aqueles games 3D mais potentes!

Possuem quatro núcleos (quad core), 8 MB de cache, e são apresentados em clocks que variam de 2.5 GHz a 3.9 GHz. São fabricados em 22 nm, assim como os Core i5.

Oferecem recursos como Intel Turbo Boost e Intel vPro, além de possuírem HT (tecnologia Hyper Threading). Não se preocupe, você saberá o que é isso daqui a pouco!

Olha aí a fera!



*Figura 2.24 – Processador Intel Core i7 (visão de cima e de baixo).*

O processador Inter Core i7 é vendido, também, em uma versão especial chamada Extreme Edition (Edição Extrema!). Traz alguns diferenciais em relação à versão “comum” (que já é muito boa), especialmente voltado para os viciados em jogos pesados. (sim, claro que é mais caro!)

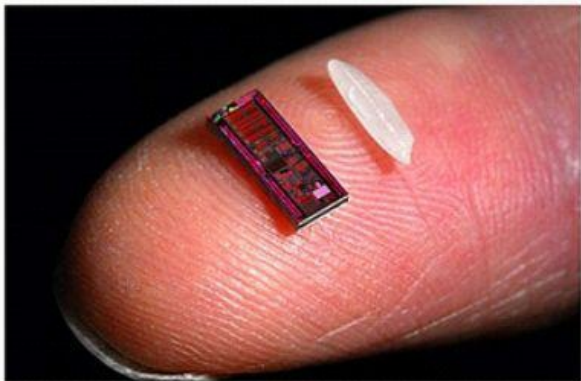
### **Intel Atom**

O Intel Atom é um processador muito pequeno e econômico, criado para o mercado de netbooks, tablets e smartphones.

Não é muito potente, mas é muito adequado ao que se destina. Seu consumo de energia é extremamente baixo, fazendo durar as baterias dos dispositivos por mais tempo. O que o torna ideal para a mobilidade.

Claro que não se pode exigir que um processador que tem quase o tamanho de um grão de arroz possa ser potente. Dá uma olhada...





*Figura 2.25 – Sim, isso daí é um dedo... E, ali no centro, o Intel Atom e um grão de arroz.*

### **Intel Xeon**

Processador criado especificamente para o mercado de servidores (computadores que ficam 24 horas por dia ligados nos centros das redes de computadores).

Com absurdos 30 MB de memória cache (alguns modelos) e até 10 núcleos (deca core), os processadores Xeon (fala-se Zíon) são designados para serviços muito além daqueles que costumamos exigir do processador (seu preço também está muito além daquilo que normalmente estamos dispostos a pagar!).



*Figura 2.26 – Processador Intel Xeon.*

Você, com certeza, pode até estar desejando um desses, mas ele não será ideal para seu uso, amigo leitor! O Xeon é feito para trabalhar como servidor, mesmo, para executar programas que têm perfil de servidor.

Apesar de muito potente e muito caro, ele não trará mais benefícios do que um bom Core i7 para a sua vida!

Há uma coisa que o Xeon pode fazer que outros processadores anteriormente apresentados não podem: Multiprocessamento. É possível montar computadores com vários processadores Xeon em paralelo.

Já imaginou? Quatro processadores Xeon com 10 núcleos cada um? São 40 núcleos no computador! É uma máquina!!!

#### **2.4.2.3. Termos do universo Intel**

Há alguns termos que descrevem tecnologias, recursos e artifícios usados somente nos processadores Intel. Os conceitos a seguir podem ser vistos em prova e foram (ou são) exclusividade da Intel.

## **MMX**

Como já foi visto, MMX é um conjunto de instruções criado para os processadores Pentium há mais de uma década.

MMX quer dizer Multimedia Extensions (Extensões para Multimídia) e consiste num conjunto com algumas instruções feitas para manipular vários dados simultaneamente. Através de tais instruções, um processador será capaz de tratar mais de um dado por vez em sua unidade lógica e aritmética.

Essa tecnologia foi essencial para transformar o computador numa máquina multimídia (para vídeo, som e imagem) e hoje em dia equipa todos os processadores da Intel.

## **HT – Hyper Threading**

O HT é uma tecnologia que, através do acréscimo de circuitos e instruções ao núcleo da CPU, faz simular a existência de dois núcleos de processamento. O HT traz ganhos reais da ordem de 20% em comparação aos processadores sem esse recurso.

Vale salientar, também, que o HT só será totalmente aproveitado se os programas usados forem compatíveis com essa tecnologia (ou seja, se o sistema operacional e os aplicativos forem escritos com as instruções específicas do Hyper Threading).

Claro que o recurso de HT é apenas um “rascunho” quando comparado à “obra de arte” do dual core (costumo dizer: “HT teve um sonho, dual core o realizou!”). Hoje, há alguns modelos de processadores com multicore + HT!

Um exemplo interessante: há alguns modelos de Core i7 que possuem quatro núcleos + HT, e, neste caso, o Windows os “enxerga” como oito processadores diferentes (é uma ilusão, já que o HT não é exatamente uma “duplicação”).

## **EM64T – Intel 64**

Se tem uma coisa na qual a Intel não foi pioneira, foi a utilização de processadores de 64 bits. A AMD começou muito antes que a Intel a explorar o mundo dos 64 bits de dados. Quando a Intel resolveu revidar, o fez criando uma série de instruções e mudanças estruturais nos chips para que estes pudessem, ao mesmo tempo, executar programas de 64 bits e manter compatibilidade com os programas de 32 bits.

A tecnologia que a Intel usa em todos os seus processadores de 64 bits é chamada de EM64T ou Intel 64 (termo mais usado agora). Se você ler, caro leitor, numa prova qualquer (ou mesmo em livros e sites da Internet), a expressão EM64T, saiba que significa, simplesmente, “processadores Intel de 64 bits”.

Vale lembrar que mesmo comprando um processador com EM64T (ou seja, um processador de 64 bits da Intel), ele só atingirá todo o seu desempenho possível quando estiver executando programas de 64 bits (especialmente o sistema operacional).

## **Execute Disable Bit – Bit de Desativação de Execução**

Este é um recurso relativamente novo na Intel e na concorrente (a AMD também tem, só que com outro nome).

Consiste em um recurso de segurança dos próprios processadores que impede que uma aplicação qualquer (como um vírus, por exemplo) possa escrever dados em áreas específicas

protegidas da memória principal. Esse recurso só será plenamente habilitado se o sistema operacional for condizente com essa tecnologia, já que é com a ajuda do sistema operacional que se definem as áreas onde se pode e onde não se pode executar instruções.

Esse recurso é bastante útil para evitar os chamados ataques de buffer overflow (sobrecarga de pilha), que é uma técnica muito usada por certos tipos de malware (programas maliciosos) para afetar a estabilidade de um sistema, culminando no seu travamento, em alguns casos extremos.

Conheça mais algumas características dos processadores da Intel no hotsite do livro na Editora Campus ([www.elsevier.com.br](http://www.elsevier.com.br)). Procure, lá, por materiais complementares a essa obra!

### **2.4.3. Processadores da família AMD**

#### **2.4.3.1. Considerações iniciais**

Embora as perguntas de provas ainda não tenham citado diferenças de arquitetura e comparações de outros critérios entre a Intel e a AMD, preterindo essa última em suas questões (sempre houve muito mais questões sobre a Intel), o aumento de participação no mercado que a AMD tem obtido nesses últimos anos tem lhe conferido um lugar de destaque que poderá ser explorado em concursos futuros.

A AMD, atualmente, é mais conhecida pela fabricação de chips de vídeo (para placas de vídeo) que de processadores. Após ter comprado uma empresa famosíssima deste segmento (a ATI), a AMD passou a assinar a fabricação das placas de vídeo Radeon.

#### **2.4.3.2. Os principais processadores da AMD**

##### **A Família Phenom**

Atualmente, a AMD desenvolve uma família de processadores multiuso chamada Phenom (“Fenômeno” é a palavra que lembra).



*Figura 2.27 – Processador AMD Phenom II.*

A família de processadores Phenom traz inúmeros modelos com diferentes Clocks e características. É usada em vários computadores pessoais Desktop e Laptop.

São famílias de produtos legados (antigos, embora ainda possivelmente vistos) da AMD os processadores Sempron e Athlon. Eles são anteriores aos Phenom, e, portanto, não merecem muita análise!

Notem: basicamente, até mesmo pela quantidade de vezes que apareceu em prova, a AMD não é muito estudada. Conheça, porém, esses nomes de processadores, para saber que pertencem à AMD!

### **Tecnologia Vision – A Nova “Onda” da AMD**

Se alguma coisa pode ser cobrada de “atualidades” sobre a AMD, é sem dúvida, a nova “geração” dela: a Tecnologia Vision.

A AMD comprou, há alguns anos, a empresa ATI, fabricante das famosas placas de vídeo Radeon (sonho de consumo de qualquer viciado em jogos de computador). As placas de vídeo, como vamos conhecer ainda adiante, são os equipamentos que constroem as imagens que são exibidas no monitor. Portanto, tudo aquilo que você vê diante de si, na tela do computador, foi construído pela placa de vídeo.

Acontece que a AMD fundiu seu antigo “mercado” com o seu novo “produto”, criando processadores (CPUs) que têm a capacidade de gerenciar, também, as imagens que o

computador constrói. Isentando o micro de necessitar de uma placa de vídeo em si.

Tal “acumulação” de funções, supostamente, traz mais velocidade aos programas e atividades que usam vídeo (ou seja, quase tudo hoje em dia!), como jogos, animações, filmes etc.

A AMD chamou essa nova “realidade” de Vision. E resolveu atribuir uma nova nomenclatura aos seus produtos, que deixaram de ser chamados de “processadores” ou “CPU” e passaram a ser conhecidos como **APU** (Unidade de Processamento Acelerado). (É... Mas só a AMD chama assim!)

Então, uma APU é, basicamente, uma “CPU” que contém, também, dentro de si, circuitos de placa de vídeo.

Atualmente a AMD fabrica esses produtos em “séries” distintas, de acordo com o nível de exigência (e potência) do computador que se deseja montar. A Série “A” traz os melhores produtos. Já as séries “E” e “C” da AMD Vision trazem processadores (desculpe, “APU”) mais simples, porém, claro, mais acessíveis!

#### **2.4.4. Palavras finais sobre processadores**

Bem, caro leitor, acho que com isso terminamos a “parte básica” sobre processadores. Ainda há algumas coisinhas a serem vistas: alguns detalhes bem pesados sobre arquitetura interna dos processadores, barramentos, entre outros, mas, veremos isso mais adiante. Vamos agora dar uma olhada na placa-mãe.

#### **2.4.5. Placa-mãe**

O processador, como foi visto, é o componente mais importante do computador, e isso não se pode negar! Como cérebro, ele tem a função de processar todas as informações que chegam a ele e devolver resultados surpreendentemente rápidos e precisos.

Mas o processador é apenas um circuito eletrônico integrado (um chip) de dimensões diminutas e corpo delicado (nossa, que texto bonito, não?). O processador tem de ser, necessariamente, ligado a uma estrutura maior e mais complexa: a placa-mãe do computador.

A placa-mãe é uma placa de circuitos, como já foi visto. Para ser mais exato, a placa-mãe é a **principal placa de circuitos de um computador**. Nela são encaixados os principais componentes de um computador, como o processador, a memória principal e os discos.



*Figura 2.28 – A placa-mãe de um computador.*

Nem é preciso mencionar que a placa-mãe tem de ser feita para o modelo específico de processador que se deseja instalar nela, não é? Ou seja, os processadores têm tamanhos e formatos de encaixe diferentes. (Esse “local” onde as CPUs são encaixadas na placa-mãe é chamado normalmente de *soquete*.)

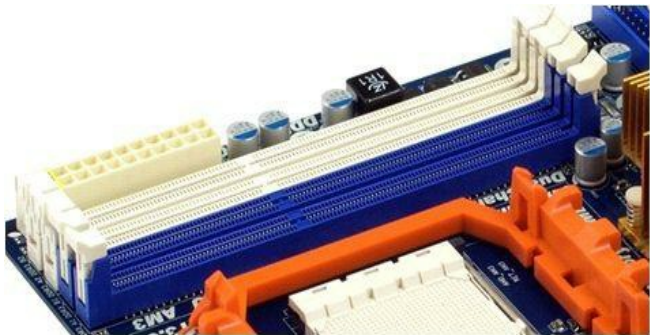
Há vários tipos de placas-mãe. Algumas são caríssimas, mas trazem recursos e desempenhos dignos de seus preços (ou talvez não), e outras são mais básicas, limitando-se a ligar os componentes do computador entre si (que, aliás, é sua função primária).

Como a placa-mãe é o equipamento em que todos os demais componentes serão encaixados, deve haver, claro, conectores (locais específicos) para que esses equipamentos sejam plugados a ela, não é? Esses conectores existem e são chamados de *slots* (fendas). De certo, slot não é um termo que serve para designar todos eles, apenas os que têm um formato de fenda, ou “rachadura” (os mais compridos e finos). Aos demais, com formatos diferentes, normalmente usa-se o termo *conector*, que é mais genérico.

Quando o conector, porém, é usado “externamente” ao gabinete (ou seja, ele pertence à placa-mãe, mas, quando montado no gabinete, aparece do lado de fora deste), chamamos normalmente de *porta* (como as “portas USB” onde ligamos nossos pen drives).

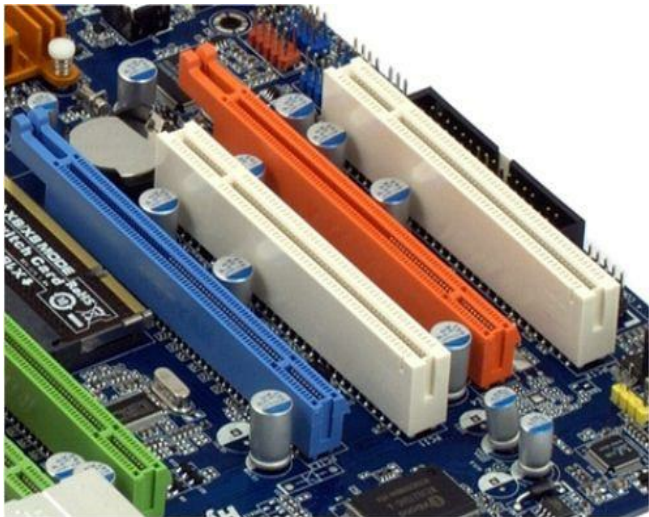
Portanto, se é “comprido e fino”, é SLOT. Se é externo, é PORTA. Se é o conector da CPU, é SOQUETE. Se não compartilha de nenhuma dessas características, chamamos genericamente de CONECTOR, mesmo!

Veja detalhes de cada um deles nas imagens a seguir:



*Figura 2.29 – Slots de memória principal.*





*Figura 2.30 – Slots para placas de expansão (como placas de som e vídeo).*





*Figura 2.32 – Soquete para a CPU.*

Continuando a lógica da coisa: os conectores não teriam a mínima função se não estivessem ligados a uma estrutura para transmitir os dados dos equipamentos a eles encaixados, não é mesmo?! É aí que entram em cena os **barramentos**. A placa-mãe, como já foi dito, é repleta de barramentos.

Barramentos são, novamente, os caminhos por onde a informação trafega entre os diversos componentes do computador. Como veremos isso mais adiante, só vou lembrar que existem duas “patentes” (ou níveis hierárquicos) de barramentos: o barramento do sistema e os barramentos de expansão (já vistos anteriormente).

Os barramentos de expansão são, precisamente, aqueles ligados aos slots (e portas) de expansão, que são os conectores usados para ligarmos os periféricos de entrada/saída e os discos.

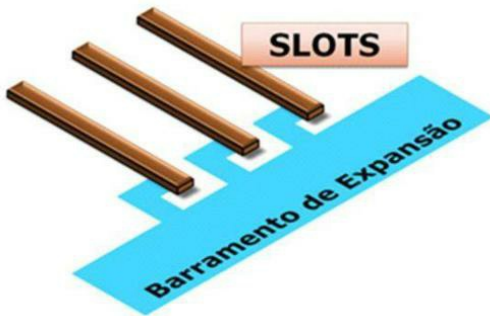


Figura 2.33 – Slots de expansão – ligados a um barramento de expansão.

Existem vários tipos de slots de expansão porque, claro, há vários tipos de barramentos de expansão. Conheceremos todos eles, um a um, mais adiante. Chegou a hora de conhecermos um componente muito importante e que já vem sendo citado há muito tempo neste livro: o *chipset*.

O chipset (ou conjunto de chips) é uma dupla de circuitos integrados presente na placa-mãe de um computador. O chipset é composto por dois chips principais: a *Northbridge* e a *Southbridge* (respectivamente *Ponte Norte* e *Ponte Sul*).

Cada um desses dois chips é composto de uma série de circuitos controladores internos para os diversos barramentos e recursos que a placa-mãe oferece. Podemos dizer que o chipset é o “cérebro” da placa-mãe, pois é seu principal componente.

Basicamente, todas as informações que trafegam entre os diversos componentes do computador têm de passar pelo chipset: ele é o “centro nervoso” da placa-mãe (e, conseqüentemente, do computador todo).



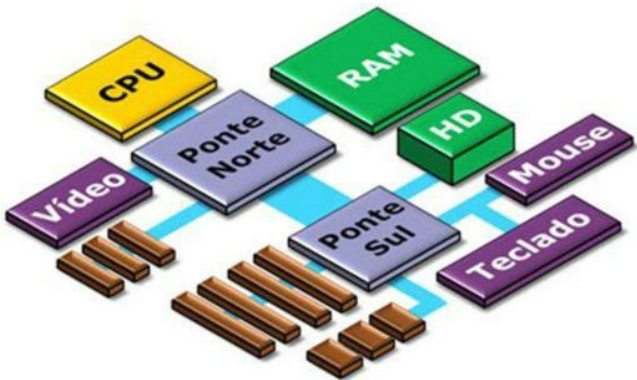
*Figura 2.34 – O chipset: Ponte Sul e Ponte Norte (a maior).*

A Ponte Norte é, sem dúvidas, o chip mais importante: é a ela que estão ligados, diretamente, a CPU (o processador), a memória principal e os barramentos de expansão mais rápidos (Atualmente, o PCI-Express).

A Ponte Sul, por sua vez, é responsável por controlar os diversos barramentos de expansão com menor velocidade, como o IDE (antigamente), o USB, o PCI e o SATA (para discos). Além, também, de possuir o circuito controlador da placa de som (caso esta seja fabricada na própria placa-mãe – isto é, caso seja onboard).

Na imagem anterior, não conseguimos VER o chipset, em si. Estamos vendo os dissipadores de calor (estruturas metálicas) instalados em cima dos chips do chipset (já que eles esquentam demais!).

Claro que veremos com mais detalhes todos esses termos descritos. Por ora, dê uma olhada no esquema a seguir que explica as ligações entre os diversos componentes do seu micro, tendo, claro, o chipset como central de comunicações.



*Figura 2.35 – Desenho esquemático dos componentes do micro.*

“Você está brincando, não é, João? Onde está o barramento de sistema? Se essa é a verdadeira estrutura atual de um micro, o que dizer da **Figura 2.12**?” (essa é a hora perfeita de voltar algumas páginas para revê-la!) – Essa é uma excelente pergunta, amigo leitor, e é muito simples de responder! **O chipset é o barramento de sistema!**

Na verdade, a ideia do barramento de sistema como uma estrada real que interliga todos os componentes do computador é apenas histórica. Hoje em dia é o chipset que faz essa interligação – ele possui, dentro de si, o barramento do sistema e os controladores dos barramentos de expansão.

Então, não se esqueça do porquê de o chipset ser considerado “a central de todas as transferências de dados no computador” – ele é o barramento de sistema! Além disso, todos os circuitos controladores dos barramentos de expansão também fazem parte dele.

#### **2.4.5.1. Controlador de memória integrado à CPU**

Atualmente, é necessário que se mencione, a Ponte Norte não é mais responsável pela comunicação entre CPU (processador) e Memória Principal (RAM). Essa conversa, na verdade, é feita DIRETAMENTE!

Ou seja, nas placas-mães e processadores atuais, compete à CPU falar diretamente com a memória principal, sem que se utilize a Ponte Norte para isso. É o que se chama de “CPU com controlador de memória integrado”. Ou seja, os circuitos que “controlam” a conversa com a memória principal estão dentro (integrados) do próprio corpo da CPU.

Desta forma, a Ponte Norte fica um tanto “limitada” a tarefas menos nobres (mas ainda assim incrivelmente rápidas), como fazer a comunicação da CPU com a placa de vídeo e com os barramentos de expansão mais velozes.

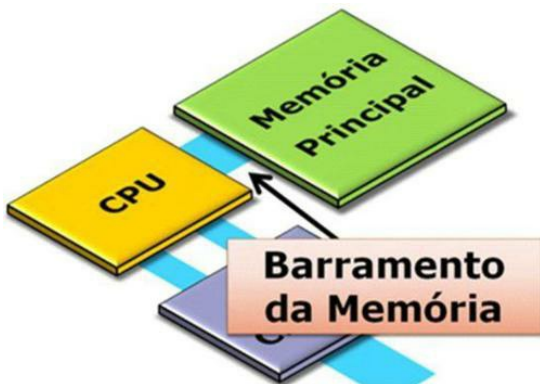


Figura 2.36 – CPU com Controlador de Memória Integrado.

Como é possível ver, ao caminho que liga a CPU à Memória Principal dá-se o nome de Barramento da Memória (um componente da placa-mãe, também!)

Com isso, amigo leitor, terminamos a análise sobre as placas-mãe e passaremos agora ao estudo das memórias de um computador.

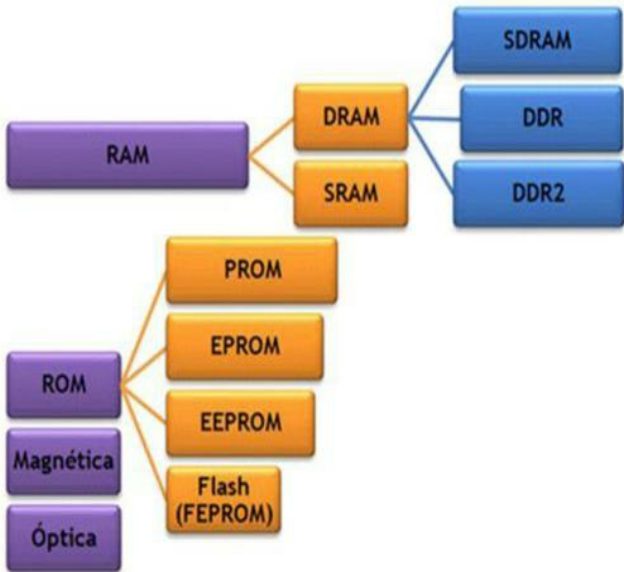
#### 2.4.6. Memórias

Como já foi visto de forma um pouco sucinta, as memórias são os *dispositivos que armazenam informações* em um computador. Existem vários tipos de memórias, desde aquelas que guardam a informação por apenas alguns instantes enquanto o computador está ligado, até as que conseguem armazenar informações por tempo indeterminado, como meses ou anos.

As memórias podem ser classificadas por seus tipos (tecnologias de fabricação) ou por suas funções (aplicações de uso no computador). Apresentarei as duas formas de classificação agora.

##### 2.4.6.1. Classificação pelo tipo de memória

Existem vários tipos de memórias, que utilizam diversas tecnologias para armazenar informações. Vamos a algumas delas:



*Figura 2.37 – Os tipos de memórias.*

### **Memória RAM**

A RAM (Random Access Memory – Memória de Acesso Aleatório) é uma memória eletrônica (ou seja, composta por circuitos eletrônicos) que armazena informações eletricamente. Por essa característica, as memórias RAM não conseguem manter os dados nelas guardados depois que o computador é desligado.

Em poucas palavras: **as memórias RAM são voláteis**, isto é, elas perdem seu conteúdo com facilidade (repito: basta que o computador se desligue para que isso aconteça). Logo, percebe-se que esse tipo de memória não foi feito para guardar informações para a posteridade, mas sim, apenas enquanto o micro estiver funcionando.

As memórias RAM podem ser, basicamente, de dois tipos:



**1. DRAM (RAM Dinâmica):** são memórias mais simples (e, com isso, são mais baratas). São feitas com circuitos baseados em **capacitores** (pequenos componentes que armazenam carga elétrica).

Como os capacitores armazenam carga, eles funcionam como baterias recarregáveis e precisam, o tempo todo, ser recarregados. As memórias DRAM precisam, portanto, de uma frequente realimentação das cargas de seus microcapacitores. Essa recarga constante é chamada **refresh** (realimentação).

O processo de refresh acontece o tempo todo nas memórias DRAM (várias vezes por segundo). Isso serve para que os capacitores com carga elétrica (que significam estado 1) não as percam (não esvaziem), o que significaria que seu estado se tornaria 0. Se qualquer bit armazenado em uma memória RAM simplesmente mudasse de estado, a confiança naquela memória seria perdida.

Chips de memória DRAM são usados em grande quantidade nas **memórias principais** dos computadores.

**2. SRAM (RAM Estática):** essas memórias são mais complexas de fabricar e, por isso, são mais caras. As memórias SRAM são muito mais velozes que as DRAM (apesar do nome) e são usadas quando a exigência de velocidade é prioridade (exemplo: nas memórias cache e nos registradores da CPU).

**PRESTE ATENÇÃO**, caro leitor! Essa “pegadinha” entre os termos “estática” e “dinâmica” pode levar você a se confundir (afinal, você sabe bem o que significam os conceitos de “estática” e “dinâmica”). Não caia nessa! A RAM estática é **mais rápida** que a RAM dinâmica!

As memórias SRAM não necessitam de refresh, pois não utilizam capacitores em sua estrutura. As memórias SRAM usam circuitos lógicos semicondutores (que consomem bem menos energia). Esses circuitos armazenam 0 e 1 por meio da variação dos estados físicos de seus componentes, e não através de cargas elétricas; portanto, não perdem as informações nelas contidas (a menos que se desligue o micro, claro).

### Subtipos da Memória DRAM

A memória DRAM já apresentou, e ainda hoje apresenta, uma série de subtipos específicos. Vamos começar com as mais antigas memórias DRAM existentes, datadas das décadas de 1970 a 1990:

- **ADRAM (Assynchronous DRAM – DRAM Assíncrona):** esse é o tipo original de DRAM, usado nos micros durante toda a década de 1970 e início de 1980 (nem precisa dizer que é antiga e não é mais usada).

**FPM RAM (Fast Page Mode RAM – RAM Modo Rápido de Página):** essa memória trouxe algumas melhorias em relação à sua precursora. Foi comum achá-la no final da década de 1980 e início da década de 1990.

- **VRAM (Video RAM – RAM de Video):** uma memória originalmente criada para placas de vídeo, por permitir maior velocidade na escrita e leitura de dados. Hoje é obsoleta.
- **EDO RAM (Extended Data Out RAM – RAM com Saída Estendida de Dados):** usada em computadores nos primeiros anos da década de 1990, oferecia mais velocidade na leitura, com ganhos de cerca de 5% em relação às memórias FPM.

- **SGRAM** (*Synchronous Graphic RAM – RAM Síncrona para Gráficos*): memória usada para placas de vídeo. Sucessora das VRAM. Essa memória também não é mais utilizada.

Atenção! Não é necessário se preocupar em decorar esses tipos que você acabou de ler, ok? Só os coloquei aqui por descargo de consciência! Isso significa que não é necessário gastar seus preciosos neurônios com esses arcaicos tipos de RAM!

(Aaahhh... A expressão é “descargo” mesmo, e não “desencargo”, antes que você critique!)

No final da década de 1990, um tipo de memória se tornou comum: A **SDRAM** (Synchronous DRAM – **DRAM Síncrona**). Esse tipo de memória possuía uma característica nova, em comparação às antigas: o fato de seu clock (frequência) ser sincronizado com a placa-mãe, ou seja, quem determinava o clock em que a memória iria trabalhar era o chipset.

A memória SDRAM evoluiu, transformando-se, anos mais tarde, naquela que conheceríamos como **DDR-SDRAM** (Double Data Rate – SDRAM, ou SDRAM com Dupla Taxa de Dados), uma memória que duplicava a transferência de Dados entre CPU e memória principal usando o mesmo clock da SDRAM.

Hoje em dia, convivemos com a memória DDR3 como mais comum para o que se usa em memória principal. Antes desta, porém, convivemos alguns anos com a DDR2, sucessora da DDR original.

Já é possível encontrar, porém, algumas placas-mãe com slot para a nova geração da memória DDR: a DDR4.

Eis, na figura abaixo, um exemplo de Pente (Módulo) de memória DDR4 da Samsung®.



Figura 2.38 – Pente de memória DDR4.

A propósito, é bom que se explique que, a rigor, a cada nova geração da memória DDR, a velocidade (taxa de transferência) duplica (em média), ou seja, para fins gerais, a DDR3 é duas vezes mais rápida que a DDR2.

### Memória ROM

As memórias ROM (Read-Only Memory – Memória Somente para Leitura) são memórias fabricadas na forma de circuitos eletrônicos integrados (chips) como as memórias RAM.

A principal diferença entre essas duas “irmãs” é que a ROM não perde o conteúdo que está gravado em seu interior, mesmo quando não há energia alimentando-a (ou seja, a ROM não é volátil como a RAM). Isso porque *os dados da ROM são gravados já na fábrica*, ou seja, a memória ROM já nasce com os dados que terá de armazenar durante toda a sua existência.

Nenhum dado pode ser alterado ou apagado da ROM. Também não é possível adicionar novos conteúdos a essa memória. Ela simplesmente poderá ter seu conteúdo lido, nunca escrito. Por esse seu funcionamento, digamos, radical, as memórias ROM parecem não ter tanta utilidade, não é mesmo, leitor?

“É verdade, João! Por que usar uma memória na qual não posso guardar meus próprios dados? Para que ela serve mesmo?” – você deve estar pensando...

Simples! Quando o fabricante de um equipamento eletrônico digital qualquer (como uma placa-mãe ou até mesmo um telefone celular) quer gravar o “comportamento” básico daquele equipamento (eu até diria a “personalidade” dele), o faz em chips de memória ROM, porque isso garante que o equipamento sempre vai funcionar segundo o que está programado em seu sistema básico (que, por estar numa ROM, é inalterável!).

Por exemplo, onde você acha que está determinado que quando você aperta o número 3 no seu telefone celular, aparecerá o 3 na tela dele? Isso está gravado num programa (o sistema operacional) no celular. Esse programa está gravado numa memória ROM no seu celular.

Esses programas básicos, ou códigos de programação escritos pelo fabricante, são comumente chamados de *firmware* (um intermediário entre software e hardware). O termo firmware descreve qualquer programa básico que determina o funcionamento de um equipamento de hardware. Tais programas são normalmente gravados em memória ROM (ou em variantes dela).

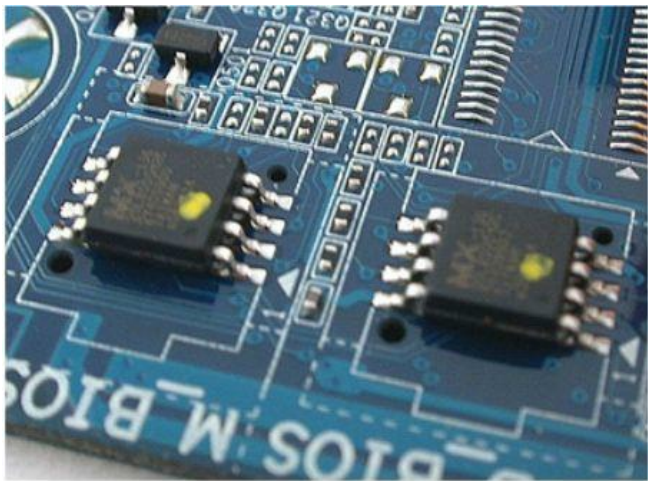


Figura 2.39 – Dois chips de memória ROM na placa-mãe de um computador.

Se analisarmos por outro ângulo, a principal característica da memória ROM é sua principal limitação. O fato de uma memória não poder ser alterada faz imaginar: o que aconteceria se

uma empresa (uma fabricante de celular, por exemplo) construiu um equipamento contendo falhas nos firmwares?

Isso seria realmente um incômodo, pois a única forma de corrigir o problema seria por meio de um recall (devolução) de todos os aparelhos com o programa defeituoso para que seus chips pudessem ser substituídos. Imagina só!

Por causa dessas e de outras, as memórias ROM deram origem a outros tipos de memórias (todas elas não-voláteis), como as que vamos conhecer agora:

- **PROM (Programable ROM – ROM Programável):** sem dúvida, essa memória pertence ao “universo” dos apaixonados por eletrônica. Ela consiste em chips que são vendidos virgens (sem dados) e que podem ser gravados apenas uma única vez. Não é muito comum vê-la sendo usada (ou mesmo citada) em textos de informática. Essas gravações acontecem em equipamentos especiais (os gravadores de PROM) e são feitas por pessoal especializado. As memórias PROM, depois de gravadas, se tornam inalteráveis como a ROM.

- **EPROM (Erasable Programable ROM – ROM Programável e Apagável):** esse tipo de memória é capaz de receber dados gravados num gravador de PROM, como sua antecessora, mas tem a vantagem de poder ser apagada caso se deseje regravá-la.

Uma memória EPROM pode ser apagada se for exposta à luz ultravioleta por certo tempo. Por causa desse sistema “estranho” de apagamento, os chips desse tipo de memória são dotados de uma janela de vidro que dá acesso ao núcleo da memória. Incidindo luz UV nessa janela por alguns minutos, o conteúdo da EPROM é completamente limpo.



Figura 2.40 – Um chip de memória EPROM – note a “janela” de vidro em cima dele.

- **EEPROM (Electrically Erasable Programmable ROM – ROM Programável e Apagável Eletricamente):** é a sucessora natural da EPROM. Pode ser apagada e gravada várias vezes, sem a necessidade de raios UV. Todo o processo de apagamento e de gravação acontece eletricamente, dentro dos chips.

A gravação e o apagamento da memória EEPROM devem ser realizados célula a célula (ou seja, de bit em bit). Não é possível apagar o conteúdo de tais memórias em blocos (vários bits simultaneamente), o que permite concluir que sua velocidade não é sua melhor característica.

Essa memória foi uma das primeiras tecnologias usadas em cartões de memória e memórias de dispositivos digitais, como máquinas fotográficas (os primeiros modelos), mas, como consumia muita energia, nunca foi vista com bons olhos para essas aplicações (as pilhas dos dispositivos descarregavam rapidamente). Hoje em dia, seu uso é muito reduzido devido à mais nova e mais “papurizada” das sucessoras da ROM.

- **Memória flash (ou FEPRM – Flash EPROM):** é uma evolução da EEPROM. Alguns autores a tratam como um subtipo da EEPROM (chamando-a de EEPROM NAND), porém é mais comum vê-la descrita como um tipo diferente da EEPROM.

As memórias flash podem ser gravadas e apagadas diversas vezes. Não há necessidade de aumento da corrente elétrica para apagá-la ou gravá-la. O processo de gravação é feito em blocos (vários bits de uma só vez), o que a torna mais rápida que a EEPROM.

Seu uso mais comum é em cartões de memória de máquinas fotográficas digitais, memórias de tocadores portáteis de MP3, as memórias dos celulares (para armazenar as agendas e compromissos) e muito mais.

Como é uma memória muito rápida e muito econômica (não em questão de custo, mas em questão de consumo de energia elétrica), a memória flash mereceu seu papel de destaque na atualidade. E ela ainda quer mais! A mais conhecida aplicação das memórias flash é, sem dúvida, os drives acopláveis às portas USB do computador, os chamados Drives Flash USB (ou pen drives).



*Figura 2.41 – Um drive Flash USB (normalmente conhecido como pen drive).*

Com isso, concluímos o estudo das memórias ROM e suas variantes. Vamos agora às memórias magnéticas em um computador.

### **Memórias Magnéticas**

Alguns dos dispositivos de memória permanente (não volátil) de um computador usam uma tecnologia antiga, mas bem-sucedida, de retenção de informação: magnetismo. Isso mesmo! Algumas memórias usam ímãs para armazenar informações.

Há muito tempo as tecnologias magnéticas são usadas para armazenar informações na forma

de campos de atração e repulsão magnética. Isso já era comum em dispositivos que não são mais tão usados hoje, como fitas cassete e fitas de videocassete e em mídias para computador, como os famosos disquetes e os ainda hegemônicos discos rígidos. Todos os equipamentos que usam memórias magnéticas são regraváveis (permitem que se gravem, apaguem e leiam as informações inúmeras vezes).

**Discos rígidos** (também conhecidos como winchesters, ou **HDs**), disquetes de vários tipos e fitas usadas em processos de backup são os mais comuns tipos de memórias magnéticas usadas na informática. Não se preocupe, caro leitor, veremos todos eles mais adiante.

### **Memórias Ópticas**

Memórias mais baratas e, por isso, muito usadas atualmente, as memórias baseadas em superfícies que refletem luz são fáceis de encontrar em diversos tipos de mídias. Os discos de CD, DVD e agora seus sucessores (os discos de HD-DVD e Blu-Ray) são feitos com essa tecnologia.

Como a tecnologia óptica de armazenamento de dados é um tanto limitada (quanto às operações que se podem realizar com tais discos), há alguns discos graváveis, mas também há discos que não podem ser regravados.

Quando chegarmos às classificações das memórias auxiliares, mostrarei tanto os diversos tipos de memórias magnéticas como as ópticas.

#### **2.4.6.2. Classificação das memórias por sua função no micro**

As memórias, em um computador, são aplicadas a certas operações específicas e, por isso, recebem nomes (ou “postos”) específicos. Aqui está a classificação das memórias em relação a suas funções dentro de um computador.





Figura 2.42 – Organização das memórias de um computador por sua função.

### Memória Principal

A memória principal, como já foi visto no início deste capítulo, é a memória em que a CPU deposita os dados e de onde ela (a CPU) lê os dados dos programas em execução. A memória principal é imprescindível para o funcionamento do micro.

Alguns autores descrevem a memória principal como tendo duas partes:

- **Memória Principal Não Volátil:** guarda os primeiros programas a serem carregados (executados) quando o micro “acorda”. Essa memória principal normalmente está presente na própria placa-mãe do computador na forma de um chip de memória ROM (ou

variantes). Seu conteúdo é um pequeno firmware (programa básico) chamado BIOS (veremos adiante). Sem BIOS, seu micro não funcionaria... Ele simplesmente não ligaria!

- **Memória Principal Volátil (ou Memória de Trabalho):** essa memória é fisicamente composta por pentes de memória RAM (aqueles SDRAM, DDR ou DDR2 que vimos anteriormente).

O BIOS, que é o primeiro programa a funcionar em um computador, está gravado num chip de ROM na placa-mãe, como vimos, mas os demais programas que usamos num computador (como o Windows e o Word) não estão na ROM!

Os programas que usamos (chamados softwares), como os sistemas operacionais e os aplicativos, são gravados em memórias auxiliares (como o HD) e ficam lá enquanto não são usados. Quando, porém, o usuário executa esses programas (coloca-os para funcionar), eles são copiados para a memória principal (a RAM) e lá ficam até o usuário desligar o computador ou fechar os programas (no X da janela).

A memória principal é, portanto, o local onde os programas têm de estar se quiserem ser executados normalmente pela CPU do computador. Se um programa está em funcionamento, está na memória principal! Costumo dizer que tudo o que você está vendo na tela do seu computador (aberto na forma de janelas) está na memória principal.

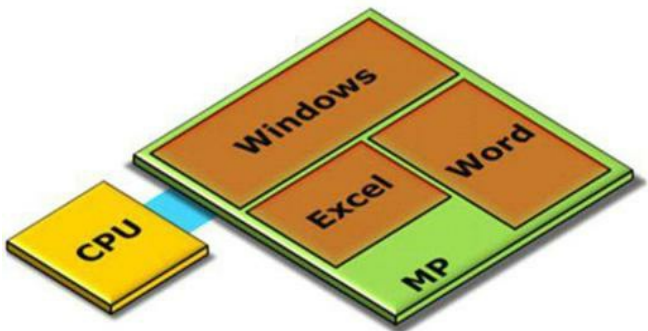


Figura 2.43 – A memória principal em ação, com vários programas sendo usados.

Lembre-se: memória RAM e memória principal não são sinônimas. Memória RAM é o **tipo** da memória. Memória principal é **função** que ela exerce em nossos computadores.

Maaaaass... Na maioria dos textos técnicos, incluindo alguns usados em provas, é comum encontrar referências à expressão Memória Principal como sinônima de Memória RAM

(acontece, portanto, você vai ter de “se ligar”).

A memória principal, e neste caso estou me referindo apenas à memória principal volátil (a RAM), tem de ter capacidade suficiente para armazenar todos os programas e dados que o usuário vai utilizar naquele exato momento. Quanto maior a capacidade da memória principal, mais programas o usuário poderá abrir sem se preocupar com o espaço.

Atualmente, os computadores podem apresentar, de acordo com as exigências que os programas fazem, de 1GB a 16GB de memória principal. O Windows 7, por exemplo, para ser utilizado em um computador que será usado apenas para texto e navegação de internet, se sente satisfeito com 2 GB (4 GB seria melhor!).

A memória principal é muito mais rápida que as memórias auxiliares (como o HD e os disquetes), mas é mais lenta que memórias como os registradores e a cache (veremos mais adiante).

## **Memória Virtual**

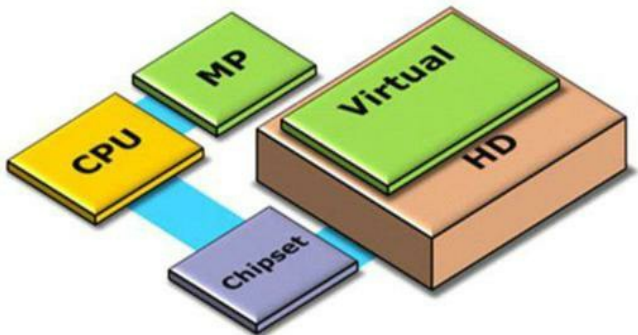
“Ei, João, o que acontece com o computador quando a memória principal está totalmente cheia e o usuário precisar abrir outro programa? O micro vai travar, não é? Ou então o programa não poderá ser aberto porque não terá espaço na memória RAM?”

Simples, leitor! O sistema operacional (no caso do exemplo, o Windows) notará que não há mais memória principal para ser usada e simplesmente cria mais memória principal!

“Como é? Agora não entendi mais nada! Explica isso direito, João!”

Com prazer! Quando não há memória principal real suficiente para abrir todos os programas que o usuário deseja, o Windows, responsável por gerenciar a MP, simplesmente “pede emprestado” um pouco de espaço da maior de todas as memórias: o HD.

Esse “espaço alugado” será interpretado pelo Windows (e por qualquer outro sistema operacional) como parte complementar da memória principal, mesmo dentro do HD. A esse recurso dá-se o nome de **memória virtual**. Se não houvesse o recurso da memória virtual, quando o usuário quisesse abrir mais programas que aqueles que a RAM consegue aguentar, provavelmente receberia uma mensagem do tipo: “Lamento, usuário. Não há espaço suficiente para abrir o programa que você solicitou. Por favor feche um ou mais programas abertos e tente novamente.”



*Figura 2.44 – A memória virtual é uma área criada no HD para complementar a MP.*

Portanto, quando um programa é maior do que o espaço livre na memória física (termo que usamos para identificar a parte real, sem ser a virtual), o sistema operacional (seja ele Windows, Linux ou qualquer outro) usará o espaço previamente reservado no HD como complemento da RAM. Se ainda há espaço na memória principal, parte do programa será copiado para lá (para ele ser executado) e o resto do programa será copiado para a área devida na memória virtual.

“João, isso é bom, não é?” – Claro que sim, leitor! Mas tem um ponto negativo.

A partir do momento em que um computador começa a usar demais a memória virtual (por estar com muitas janelas abertas ou por ter pouca memória RAM instalada, ou quem sabe ambos), o computador começa a deixar de trabalhar com uma memória rápida (a RAM física) e passa a usar uma memória dezenas de vezes mais lenta (o HD, onde a virtual é criada). Portanto, tenha muita memória RAM para não precisar depender da virtual, pois se isso acontecer, seu micro virará uma “carroça”.

Em resumo, lembre-se: é bom ter muita memória RAM instalada. Isso fará com que o micro não precise usar a virtual e, com isso, não caia de desempenho. É por isso que os vendedores vivem dizendo: “compre muita memória RAM, pois quanto mais memória for instalada, mais rápido o micro será!” (Sabemos que não está totalmente certo, mas tem um fundo de verdade!).

Em tempo, a frase que os vendedores deveriam dizer é: “compre muita memória RAM, porque se seu micro tiver pouca, ele vai precisar usar mais frequentemente a memória virtual e com isso ficará mais lento!”.

Existem várias técnicas de se fazer memória virtual. Cada sistema operacional usa a que foi programado para usar. As três técnicas usadas para fazer memória virtual são a paginação, a segmentação e a troca (swapping).

- **Paginação:** a memória principal é dividida em pequenos blocos chamados páginas. As páginas têm sempre o mesmo tamanho (ou seja, os tamanhos das páginas são independentes do tamanho dos programas a serem executados na memória).

Contrariando o que ainda vamos ver (a execução de um programa), no caso das memórias virtuais paginadas, um programa poderá ser gravado em áreas não contíguas da memória, ou seja, um programa poderá ser gravado “espalhado” por diversas páginas diferentes em posições diferentes da memória principal.

Graças à divisão da memória principal (considerando-a toda: física + virtual) em pedaços, um programa poderá estar gravado parte na memória física e parte no disco (virtual).

- **Segmentação:** nesta técnica, a memória principal (física + virtual) é dividida em blocos chamados segmentos. A principal diferença entre os segmentos e as páginas é que eles têm tamanho variável. Isso significa que programas maiores (com mais instruções) são armazenados em segmentos maiores na memória.

Como há divisão da memória, da mesma forma como na memória paginada, a memória segmentada permite que um programa seja gravado de forma não contígua (espalhado) pela memória. E, claro, de forma idêntica à paginada, as memórias virtuais que usam segmentação permitem que partes do programa estejam gravadas na memória física enquanto outras partes ficarão no disco.

- **Troca (Swapping):** as memórias que fazem swap (troca) funcionam de forma muito mais “arcaica” que as técnicas que dividem a memória em pedaços. Em uma memória virtual de troca, um programa não pode estar simultaneamente no disco e na memória física. Ou está em um, ou no outro.

“Eita, João, e como isso pode ser memória virtual, hein?”

Fácil, caro leitor! Quando um programa estiver no disco e for necessário, ele é totalmente trazido para a memória física (RAM), num processo que chamamos de swap in. E, claro, se há na memória RAM alguém (algum programa) que já não está sendo tão requisitado, para dar lugar aos programas que vêm, ele terá de ser levado (totalmente) para o disco (swap out).

Nota-se a incrível “dificuldade” que se apresentará quando algum programa precisar fazer a troca, não é mesmo? Pense, leitor: tem momentos nos quais seu computador simplesmente “para”, e mostra aquela ampulheta (relógio de areia) na sua tela. Adivinha? Esse é o momento da troca! Quando seu computador precisa interromper o processamento momentaneamente para fazer a troca (swap).

Portanto, ao ler, em alguma prova, a expressão memória virtual paginada, arquivo de troca (swap file), memória segmentada, memória de paginação etc. não se espante, é apenas a boa e velha memória virtual que você acabou de conhecer!

## **Memórias Auxiliares (Memória Secundária)**

Um computador também é composto de um conjunto de equipamentos que permitem o armazenamento de informações permanentemente. Esse grupo de dispositivos (alguns com memória magnética, outros com armazenamento óptico e alguns com memória flash) é conhecido, normalmente, como memória secundária ou memória auxiliar.

Nas memórias auxiliares, as informações (programas ou dados do usuário) são armazenadas na forma de blocos de dados chamados arquivos. Os arquivos são identificados por nomes

específicos e são listados em compartimentos chamados diretórios (pastas).

Em suma, ***nós salvamos nossas informações*** justamente nas memórias auxiliares.

Vamos aos mais importantes:

#### *Disco Rígido (HD ou Winchester)*

Das memórias auxiliares, sem dúvida essa é a mais importante. Todos os computadores possuem uma memória magnética de grande capacidade para armazenar todos os programas e arquivos do usuário.

Discos rígidos são formados por vários discos metálicos sobrepostos, que giram ao redor de um eixo e são lidos (e gravados) por pequenos dispositivos magnéticos (chamados cabeças de leitura/gravação) que ficam na ponta de braços que se movem das proximidades do centro do disco para a sua extremidade.



*Figura 2.45 – Um HD aberto (a essa altura, não serve mais para nada!).*

Esses discos metálicos são magnetizáveis, ou seja, podem receber influência dos campos magnéticos gerados nas cabeças de leitura/gravação. Com o estímulo certo nas cabeças de leitura/gravação, os discos reorganizam os pequenos ímãs em sua superfície para que se posicionem no intuito de fazer significar 0 (zero) ou 1 (um).

Discos rígidos são muito velozes se comparados a outros dispositivos de memória secundária

(como disquetes, CDs e DVDs), mas são dezenas de vezes mais lentos que a memória principal. Isso se deve ao fato de o disco rígido precisar de um processo mecânico de acesso aos dados gravados em seus pratos (os discos metálicos) – ou seja, é necessário que peçam que se movimentem dentro do corpo do HD.

Para aumentar a velocidade de acesso aos dados neles contidos, os discos rígidos possuem uma memória cache! Não! Não é a mesma memória cache do processador.

A memória cache dos discos rígidos é uma pequena quantidade de memória DRAM (normalmente 32 MB, mas já há discos com 64 MB e até 128 MB) que armazena conjuntos de dados recentemente acessados para que, quando forem requisitados novamente, não seja necessário buscá-los nos discos, girando-os mecanicamente. A *cache de disco* (nome usado comumente para ela) torna o acesso aos dados dos discos muito mais rápido. Quanto mais cache de disco, mais rápido será o HD.

Atualmente, porém, temos outro “xodó” para a função que os discos rígidos executam: aqui está um exemplar de um disco rígido (não seria certo continuar chamando-o de “disco”) todo feito com memória flash (FEPROM). Esse equipamento é muito mais veloz e silencioso que um disco rígido convencional, pois não tem partes mecânicas móveis (o acesso é todo feito eletronicamente, como qualquer memória flash). Em suma, ele é como um “grande pen drive” fixo dentro do seu computador.



*Figura 2.46 – SSD (HD feito de memória flash) da Toshiba® para laptops.*

Além das vantagens citadas anteriormente, os HDs flash consomem muito menos energia que os HDs magnéticos. Um HD flash, só para se ter ideia, consome cerca de 30% da energia de um HD convencional.

A questão negativa é o custo dessa nova tecnologia. Um HD flash é muito caro se compararmos sua capacidade. Enquanto já é possível encontrar facilmente HDs magnéticos de 1 TB e 2 TB (terabytes), os discos flash mais comuns no mercado têm entre 128 MB e 512 MB (megabytes), sendo estes últimos, muito mais raros (além de caríssimos)!

Além disso, claro, o custo de um HD de 2 TB é menor do que um SSD de 256 GB!

Só para situá-lo, leitor, os “HDs” Flash são chamados de **SSD** (Solid State Disks – Discos de Estado Sólido). Este é, portanto, o termo que deverá ser usado em provas para descrever tais equipamentos.

*Disquete de 3 ½ Polegadas (disquete convencional)*

Por favor, leitor: um minuto de silêncio pela “morte” do disquete!



...

Obrigado!

O famoso disquete de 3 ½ polegadas é um disco feito de material plástico bastante flexível (como os filmes usados nas fitas cassete), envolto numa capa de plástico rígido. O disquete foi, durante muito tempo, a mídia removível mais usada no mundo da informática.

O disco plástico onde as informações são gravadas é magnético (ou seja, grava informações na forma de alterações de estados magnéticos dos pequenos componentes do disco). Hoje em dia, sua capacidade e sua velocidade são dignas de pena! É realmente triste ver algo que nos serviu tão bravamente ser suplantado, rapidamente, por dispositivos menores (em tamanho), porém bem maiores (em capacidade e velocidade).

E viva o pen drive! Onde será o “velório” do disquete? (Para eu fazer uma festa!)



*Figura 2.47 – Disquetes de 3 ½ polegadas.*

Os disquetes de 3 ½ polegadas têm capacidade “oficial” de armazenamento de 1,44 MB (megabytes). Porém, quando formatados (preparados para uso) no Windows, apresentam apenas 1,37 MB de capacidade útil.

Normalmente, o valor exigido pelas bancas examinadoras é o primeiro (o “oficial”), mas quando o Cespe/UnB resolve perguntar de forma mais prática sobre os processos de cópias de arquivos para dentro dos disquetes, deve-se considerar seu tamanho real, 1,37 MB.

#### *CD (Compact Disk)*

O CD é um disco plástico que possui uma superfície capaz de refletir a luz. E é justamente essa superfície que armazena os dados. O equipamento que lê o CD (conhecido como CD player ou

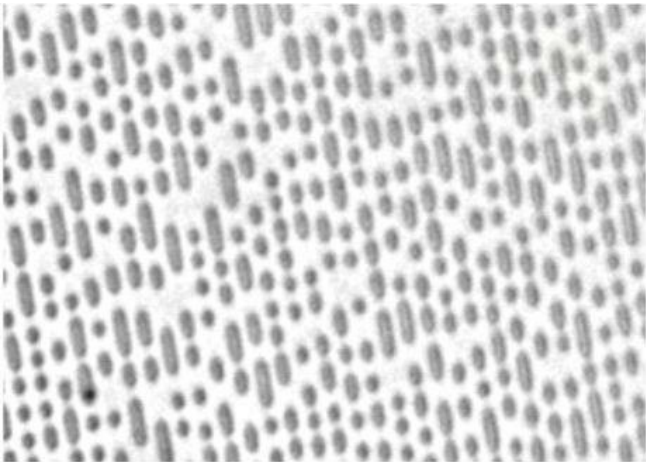
drive de CD) possui um canhão que dispara um feixe fino de laser que deverá ser refletido ao equipamento por essa superfície legível.

O CD tem, normalmente, entre 650 e 700 MB, o que equivale aos dados de cerca de 500 disquetes (pouco mais de 480, para ser mais exato).



*Figura 2.48 – Um disco de CD-ROM.*

A superfície dos CDs é formada por alguns “buracos” chamados pits (poços), que são feitos em baixo relevo na superfície mais alta, chamada land (solo). A distribuição desses pits e a forma como eles são “cravados” na estrutura do CD representam os 0 (zeros) e 1 (uns) das informações gravadas neles.



*Figura 2.49 – Dois detalhes da estrutura física de um CD (os pits são os buracos escuros).*

Existem alguns tipos de CDs usados normalmente para computadores:

- **CD-ROM:** é o CD que já sai de fábrica com dados gravados. Esse CD não poderá ter seu conteúdo alterado pelo usuário. Ou seja, nem pense em apagar, alterar ou adicionar informações nesse CD! Sua estrutura física (pits e lands) é gravada industrialmente como uma escultura na superfície de vidro; portanto, impossível de alterar.
- **CD-R (CD Gravável):** esse CD possui uma camada fina de resina em sua superfície gravável. Essa camada fina de resina será “queimada” pelo laser do equipamento gravador de CD (desses que você usa no computador mesmo). Essa “queima” cria áreas com características diferentes de reflexão do laser (o que imita os pits e lands). Os CD-R não têm pits e lands, mas possuem áreas que refletem a luz de forma diferente entre si, imitando o comportamento do laser quando lê os pits e lands do CD-ROM.

Vale salientar, também, que um equipamento gravador de CD tem dois tipos de raios laser: o mais forte serve para gravar (queimar a resina) e o mais fraco serve apenas para ler o CD (é o mesmo raio usado para ler o CD-ROM).

Claro que, depois de queimada uma área do CD-R, ela não poderá ser queimada novamente; portanto, uma vez gravado um dado no CD-R, ele não poderá ser apagado.

Mas lembre-se: um CD-R pode ser gravado várias vezes.

“Agora lascou tudo de vez... João, como é que o CD-R *não* pode ser apagado, mas *pode* ser gravado várias vezes?”

É simples, amigo leitor... O entendimento de “gravar várias vezes”, para as bancas que já fizeram perguntas nesse sentido é de que um CD-R pode ser gravado por partes – cada “parte” seria uma gravação diferente. E é assim mesmo! Quando se grava um CD-R, não se é obrigado a gravá-lo por completo: podemos gravar uma pequena parte (digamos, uns 10%) e, depois disso, gravar os outros 90% em outras oportunidades, sempre somando ao que já se tinha.

Então, o entendimento de “gravar várias vezes” é relacionado a “gravar à prestação” e não a “gravar por cima do que já se tinha anteriormente”. Essa gravação à prestação é chamada multisessão.

É multisessão porque cada “prestação”, ou seja, cada área de gravação ininterrupta, é chamada de sessão. Veremos mais adiante que a gravação de uma sessão tem de ser ininterrupta.

• **CD-RW (CD Regravável):** esse tipo de CD pode ser gravado e apagado diversas vezes (segundo os fabricantes, mais de mil vezes).

Esses CDs utilizam uma mistura de componentes químicos (prata, telúrio, antimônio e índio, para ser mais exato) em sua superfície gravável que, em seu estado normal, é cristalina sólida (consegue refletir o laser leitor) e que simplesmente se torna líquida quando aquecida pelo laser do equipamento gravador (é o momento do “apagamento” do CD-RW).

Enquanto a mistura continuar nesse estado “amorfo” (sem forma, líquido), o CD-RW não consegue refletir a luz que incide sobre ele (ou seja, o leitor vai ler “tudo vazio” no CD). Quando essa mistura esfria, ela se torna cristalina novamente e vão se formando áreas com diferentes capacidades de reflexão do laser (definidas pelo laser gravador). Essas áreas parecem os pits e lands do CD-ROM.

Vê-se a complexidade desse processo, não é? O CD-RW não pode ser apagado parcialmente (como os discos magnéticos ou memórias flash). O CD-RW só pode ser apagado em sua totalidade (formatação).

“Ahhh! Peraí, João... Essa eu não engulo não... Já copiei arquivos para dentro de um CD-RW e eles substituíram os arquivos anteriores com mesmo nome! Para onde foram os outros arquivos? Isso é uma prova de que eu consegui apagar somente aqueles arquivos e não o restante do CD-RW, como você acabou de falar.”

Pois é, caro leitor... Os programas gravadores de CD e DVD querem, a todo custo, fazer as memórias ópticas serem tão simples de gravar e desgravar como as memórias magnéticas ou as flash. Mas isso não é possível.

O seu “arquivo anterior” ainda está gravado no CD (seja R ou RW). O “novo arquivo” foi gravado em uma sessão posterior (uma nova “prestação”) e foi criada uma nova tabela de alocação no CD que aponta para o novo arquivo e simplesmente ignora a existência do anterior. Ou seja, os dois arquivos existem e ocupam espaço no seu CD. Você só acessará o mais novo deles, dando-lhe a impressão de que realmente este substituiu o anterior. Lastimável, não?

Outra coisa interessante de saber é que os CDs (de qualquer tipo) são gravados em espiral, e não em círculos concêntricos, como os discos magnéticos. Por isso, é preferível que se tenha uma elevada quantidade de dados para gravar de uma única vez que gravar diversas vezes.

Agora vamos falar um pouco dos equipamentos que trabalham com CDs:

- **Leitor de CD (Drive de CD):** é um equipamento que não é mais vendido. O leitor de CD permite apenas que um CD com dados (ou música) seja lido. Claro que não será possível gravar qualquer tipo de dados em CDs através deste tipo de equipamento.

A velocidade de leitura de CDs de música (que desde a década de 1980 continua a mesma) é de 150 KB/s (150 Kilobytes por segundo). Isso significa que, ao tocar um CD de áudio normal, o CD player trará, a cada segundo, 150 KB. Essa foi a velocidade ideal encontrada para a leitura de música.

Como os discos de dados (CD-ROM e afins) não contêm música, mas dados de computador, a velocidade de leitura pode ser muito maior que essa. É por isso que hoje há leitores de CD para computador a 60x (que significa 60 x 150 KB/s – ou 9.000 KB/s). Portanto, o “x” dos equipamentos de CD é equivalente, já deu pra perceber, a 150 KB/s.

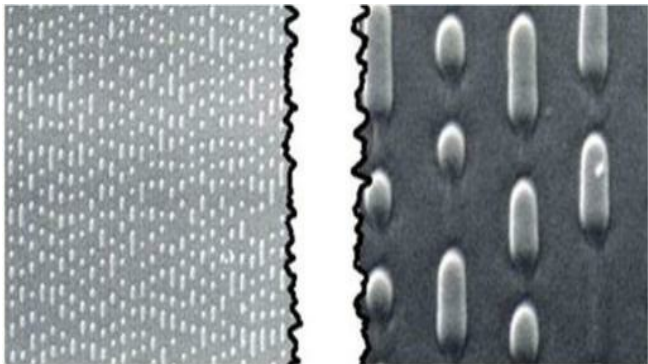
- **Gravador de CD:** os gravadores de CD conseguem, além de ler CDs de vários tipos, gravar CD-R e CD-RW e, claro, apagar CD-RW. A velocidade de gravação de CD-R é sempre menor que a velocidade de leitura, e a velocidade de gravação de CD-RW sempre é menor que a de gravação dos CD-R.

Atualmente, os gravadores de CD e DVD possuem sistemas de proteção para evitar perdas de dados. Um dos principais problemas relacionados à gravação de discos ópticos é o Buffer Underrun, que consiste no gravador não possuir os dados necessários para fazer a gravação (a memória temporária do gravador, chamada Buffer, esvazia e fica sem dados para gravar no disco). Isso faz com que o gravador passe a queimar o CD sem nenhum dado, ou simplesmente pare de queimar e, em ambos os casos, isso pode inutilizar o disco.

### *DVD – Digital Versatile Disk*

O DVD é um disco óptico, como o CD, porém, com uma capacidade de armazenamento sete vezes maior. Um DVD pode armazenar cerca de 4,7 GB (4,7 Gigabytes) de dados. Essa, vale salientar, é a capacidade de um DVD normal, de camada simples, visto que já existem os DVDs de camada dupla, com 8,5 GB de capacidade.

Fisicamente, os DVDs são muito parecidos com os CDs (em diâmetro e espessura). Mas os pits e lands do DVD são mais próximos e bem menores. Eles podem ser assim porque os lasers dos equipamentos de DVD são mais finos que os lasers disparados pelos equipamentos de CD.



*Figura 2.50 – Pits do DVD x pits do CD (a diferença é grande, não?).*

Claro que um equipamento que lê DVD consegue ler perfeitamente um CD, mas o contrário não é verdade. Equipamentos de CD são incapazes de ler discos de DVD.

As velocidades de leitura de um disco de DVD também podem variar de equipamento para equipamento. Já há leitores de DVD de 20x.

“Só isso, João? 20x? Que miséria! Os leitores de CD chegam a 60x.”

Calma, nobre leitor! O “x” do DVD é diferente! Ele se refere à velocidade de leitura de DVDs de filme, que é de 1.321 KB/s, ou seja, 1,3 MB/s (cerca de 9 vezes a velocidade de 150 KB/s dos CDs). Então, um “x” de DVD vale 9 vezes um “x” do CD.

Então, leitores de DVD com 20x são equivalentes, em velocidade, a supostos leitores de CD de 180x. Um absurdo, não acha? Ainda há velocidades diferentes para os processos de leitura, gravação e regravação dos diversos tipos de DVD. Mas, quais são esses diversos tipos?

- **DVD-ROM:** são os DVDs de filme e de programas de computador que já vêm de fábrica gravados com dados. O usuário não poderá alterar-lhes nenhuma característica.
- **DVD-R, DVD+R:** são os DVDs semelhantes ao CD-R. Ou seja, eles podem ser gravados várias vezes (em multissessão), mas não podem ser apagados.
- **DVD-RW, DVD+RW:** são semelhantes ao CD-RW. Podem ser gravados e apagados diversas vezes, mas apagados só em sua totalidade. Lembrem-se: não dá para apagar apenas uma parte destes discos.

“Certo, João... Mas qual a diferença entre -R e +R? E -RW e +RW?”

Briga da indústria, caro leitor! Alguns fabricantes apostam que os formatos -R e -RW são ideais. Outros fabricantes, porém, para não dar o braço a torcer (e para ganhar um pouco com

“royalties”), decidiram investir na família “+”!

Lembre-se de que um gravador de DVD-R (e -RW) não pode gravar (pelo menos não oficialmente) DVDs dos tipos +R e +RW. Os gravadores da família “+” também não podem gravar discos de DVD-R e DVD-RW.

Claro, caro leitor, que há inúmeros equipamentos “toca tudo” e “grava tudo” à venda! Logo, já é possível encontrar gravadores de DVD que admitem gravações nas duas “facções” (+ e -).

- **DVD-RAM:** são discos de DVD criados, primordialmente, para o mercado de vídeo (filmadoras digitais) e para os gravadores de DVD domésticos, que gravam conteúdo da TV. Esse DVD é regravável (inclusive parcialmente), o que o torna um forte concorrente dos DVD-RW e DVD+RW. O grande problema é o custo desta mídia, normalmente superior aos anteriores. Vale salientar, também, que nem todo gravador de DVD para computador é compatível com DVD-RAM, fato que ainda dificulta sua popularização no mercado de informática.

Ao contrário de todos os demais tipos de DVD e CD, que gravam dados de forma espiral (uma única linha que começa no meio do disco e vai até a borda), os DVD-RAM usam círculos concêntricos, de forma bem semelhante ao que acontece nos HDs e disquetes.

- **DVD+R DL:** também conhecido apenas como DVD DL ou DVD9, esse disco é composto por duas camadas de gravação sobrepostas numa mesma face (num único lado). Os DVD DL (DL vem de “Dual Layer” – “Dupla Camada”) possuem uma capacidade de 8,5 GB (quase o dobro da capacidade dos DVDs normais).

Para a gravação de discos de DVD DL, é necessário que o gravador seja específico para essa operação (atualmente, todo gravador de DVD, praticamente, já é DL).

Quanto aos equipamentos usados para leitura e gravação de DVDs, podem-se destacar:

- **Leitor de DVD (Drive de DVD):** é o equipamento que consegue apenas ler discos de DVD (ou seja, não consegue gravá-los). Esse equipamento também consegue ler discos de CD.
- **Drive Combo:** é um equipamento “intermediário” que consegue ler e gravar discos de CD, mas só consegue ler discos de DVD (não os grava). O termo “combo” vem de “combinação”, pois esse equipamento é uma combinação entre gravador de CD e leitor de DVD. Durante muito tempo esse foi o dispositivo para discos ópticos dos laptops no mercado (os laptops mais baratos ainda vêm com ele).
- **Gravador de DVD:** esse é o “completo de tudo”... O gravador de DVD consegue ler e gravar discos de CD e DVD. Vale apenas lembrar que esses equipamentos gravam, normalmente, apenas uma família de DVDs (ou +, ou -).

### *BD (Blu-ray Disc)*

A novíssima geração dos discos ópticos conta com um integrante de peso, sucessor do DVD: o BD, ou Blu-ray Disc, ou, simplesmente, **Blu-ray!**

Um disco Blu-ray é um disco óptico, com capacidades de até 54 GB (se for de duas camadas) ou até 27 GB (se for de camada simples).

Obviamente, só é possível ler um BD se o computador possuir um equipamento específico para esta tecnologia (drive de BD, ou leitor de BD). E, também, lógico, gravar um BD é tarefa, unicamente, para equipamentos **gravadores de Blu-ray!**



Figura 2.51 – Disco de BD-R da Sony®.

Os tipos de discos de BD são três:

- **BD-ROM:** já vem gravado de fábrica, seu conteúdo só pode ser lido. Não pode ser gravado nem apagado.
- **BD-R:** pode ser gravado várias vezes. Não pode ser apagado.
- **BD-RE:** é o BD regravável. Pode ser gravado e apagado várias vezes.

#### *Memória Flash USB (Pen drive)*

Durante muito tempo, achou-se que os CD-RW seriam os substitutos do disquete convencional por sua capacidade de armazenamento e pela característica de serem regraváveis, embora toscamente. Pois estávamos todos enganados! Eis o disquete da atualidade (e do futuro): um dispositivo de memória flash que pode ser acoplado a qualquer porta USB no computador – o pen drive.





*Figura 2.52 – Um pen drive (ou dispositivo de memória flash USB) da Kingston®.*

Já se podem encontrar facilmente pen drives com até 128 GB de capacidade (isso vai aumentar), mas os mais comuns são os de capacidades em torno de 4 GB a 16 GB. Eles recebem o nome de pen drive (ou “drive caneta”) por causa de seu formato característico (nas primeiras gerações) com tampinha para encaixar no bolso.

Por serem incrivelmente práticos (são carregados como chaveiros, não é mesmo?) e possuírem as vantagens de ser memória flash (regraváveis em blocos, pouco consumo de energia, memória não volátil de qualidade, boa velocidade de gravação e leitura) além de serem encaixados em qualquer micro, os pen drives são realmente uma “mão na roda” para transportar dados.

#### *Fitas Magnéticas (Fitas para Backup)*

Embora não sejam muito comuns em nossos computadores, as fitas magnéticas são muito usadas em ambientes corporativos, para os quais a realização de backups (cópias de segurança) é imprescindível e deve ser feita constantemente.

Como o nome já diz, essas fitas (cartuchos plásticos que contêm fitas enroladas, como as fitas de videocassete) são magnéticas e, por isso, totalmente regraváveis. (Lembre-se de que tudo que for magnético é regravável!).

“João, essas são aquelas fitas chamadas Fitas DAT?” – você pergunta...

Sim! E não! Fita DAT é apenas um dos modelos de fitas usado no mercado. Ou seja, “fita magnética” é um gênero, “DAT” é uma espécie (um subtipo).



*Figura 2.53 – Um exemplo de uma fita DAT-72 da HP®.*

As fitas atuais podem chegar a dezenas de gigabytes de capacidade, como as fitas do tipo DAT-160, que armazenam até 80 GB e suas sucessoras, as DAT-320, que conseguem armazenar centenas de gigabytes.

Essas fitas não são usadas para transporte de dados. Elas são usadas, até pela dificuldade de gravação e leitura, para backups. Ou seja, nestas fitas são guardados dados que se julgam importantes. Se alguma coisa acontecer com esses dados no local original onde se encontram, recupera-se o conteúdo deles contido nessas fitas.

#### **2.4.7. Dispositivos de entrada e saída**

São considerados dispositivos de entrada e saída, como já se viu anteriormente, aqueles equipamentos que permitem a entrada e a saída de dados da CPU do computador. Genericamente, são equipamentos que permitem que o usuário “converse” com o micro e vice-versa. É comum utilizar o termo “periféricos” para nos referirmos a eles.

##### **2.4.7.1. Teclado (entrada)**

É o equipamento que permite a inserção de dados através da digitação. É conhecido como *periférico padrão de entrada*.



*Figura 2.54 – O teclado é considerado o equipamento padrão de entrada.*

O teclado que usamos atualmente (na verdade, que sempre foi usado, desde a época das máquinas de datilografia) é chamado QWERTY em alusão à distribuição das primeiras letras na primeira linha do teclado. Pesquisas indicaram (pelo menos é o que se diz) que essa disposição (lay out) das teclas é a mais “eficiente” de todas (para a língua inglesa, claro).

Apesar de ser um padrão internacional, a disposição das teclas do teclado em QWERTY é um pouco diferente de país para país, portanto, é possível que haja incompatibilidades entre certos programas e teclados.

Aqui no Brasil, por exemplo, os teclados que apresentam a tecla de Ç (cedilha) são chamados teclados ABNT (por seguirem as normas descritas pela Associação Brasileira de Normas Técnicas). Além dessa tecla em específico, mudam as posições de alguns acentos (agudo, circunflexo, til etc.) em comparação com os teclados americanos. Portanto, apesar de, em nível mundial, usarmos teclados QWERTY, eles não são todos exatamente iguais (pequenas diferenças em relação a acentos e caracteres especiais de certos idiomas).

Além do padrão QWERTY, que é, sem dúvida alguma, o “rei do pedaço”, há outro layout muito diferente conhecido como DVORAK. Os teclados DVORAK diferem dos QWERTY na posição das próprias letras do alfabeto. Veja um exemplo a seguir.



*Figura 2.55 – O teclado Dvorak (estranho, não?).*

Os teclados DVORAK definitivamente não são padrão. É muito raro encontrar esse tipo de teclado no Brasil (mas, sim, é possível). Se você quer ser “diferente” de todo mundo, e, segundo especialistas e entusiastas, se quiser ter mais eficiência ao digitar, esse talvez seja o seu teclado ideal.

#### **2.4.7.2. Monitor (saída)**

O monitor de vídeo é considerado o periférico padrão de saída, ou seja, a saída de dados acontece preferencialmente neste equipamento.



*Figura 2.56 – Monitor de vídeo LCD (cristal líquido).*

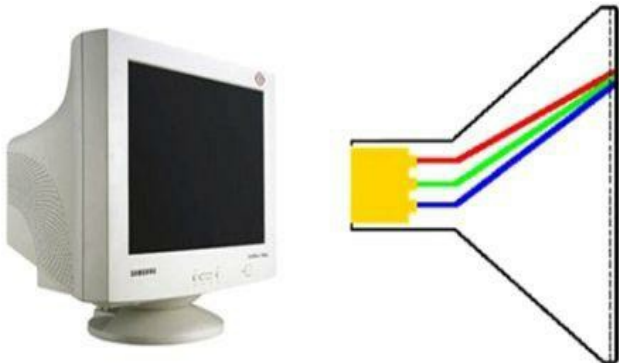
A principal característica de um monitor é o tamanho de sua tela, medido de forma semelhante à televisão. Atualmente são comuns os monitores de 15, 17, 19, 20, 23 e até 30 polegadas). Essa medida refere-se à diagonal da tela do monitor.

Os monitores podem ser classificados como monitores de **CRT** (tubos de raios catódicos) e **LCD** (monitores de cristal líquido). Os primeiros são aqueles que normalmente temos em casa, monitores volumosos, que usam um canhão que dispara feixes magnéticos numa malha de fósforo. Os monitores de LCD são finos, e normalmente são encontrados em computadores portáteis (mas também comuns em computadores de mesa).

#### **Monitor de CRT (tubos de raios catódicos)**

Um monitor de CRT emite sinais luminosos graças ao disparo de raios eletromagnéticos (os tais “raios catódicos”) originados de um canhão na parte traseira do equipamento. Esse “canhão” é conhecido como emissor de elétrons (ou arma de elétrons).

Os monitores de CRT não são mais encontrados atualmente no mercado, tendo sido quase que completamente substituídos pelos monitores de LCD.



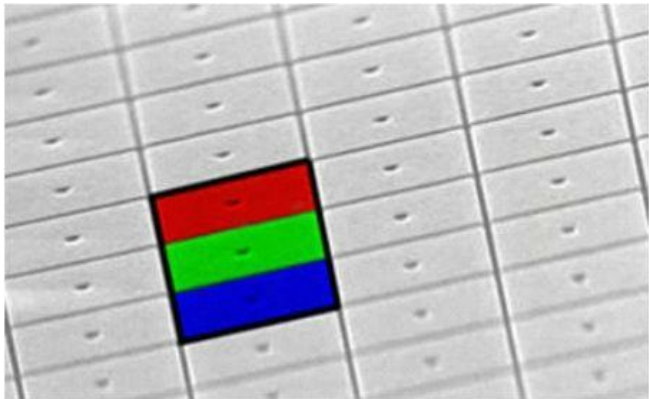
*Figura 2.57 – Monitor CRT e os três feixes eletromagnéticos disparados.*

Nos monitores coloridos (ou seja, todos atualmente), três feixes eletromagnéticos são disparados: um vermelho, um verde e um azul. Essas três cores são consideradas as cores primárias em um monitor e são suficientes para, misturadas, exibirem qualquer cor que conseguimos enxergar num monitor.

### **Monitores de LCD (tela de cristal líquido)**

Monitores de cristal líquido são conhecidos como os “monitores fininhos”. Eles constroem a imagem por meio de células retangulares na tela que deixam a luz passar quando recebem sinais elétricos. Essas células são compostas de material líquido que se cristaliza quando recebe alimentação elétrica.

As células são dispostas de três em três (nas cores primárias: vermelho, verde e azul) e juntas formam o que chamamos de pixel (embora esse termo seja usado para definir também cada quadradinho que forma a imagem digital no computador).



*Figura 2.58 – Malha de células de um monitor LCD (e um pixel – três cores – destacado).*

Os monitores de LCD são construídos em duas tecnologias diferentes: matriz passiva (antigos) e matriz ativa (atuais). Os monitores de matriz passiva eram lentos (para redesenhar os pixels), tinham pouco contraste (era difícil identificar detalhes na imagem) e apresentavam um ângulo limitado de visão (normalmente, a imagem só era percebida com perfeição quando se olhava de frente para ele).

Todos os monitores de laptops e monitores de LCD para desktops usados atualmente são construídos em matriz ativa (chamados de TFT – Thin Film Transistor – Transistor de Película Fina). Esses monitores, além de apresentarem uma qualidade de imagem superior (maior contraste e melhores cores), são mais rápidos (menor tempo de resposta e redesenho da tela) e apresentam ângulos de visão mais amplos (permitem ver a imagem do monitor mesmo quando não se olha diretamente para ele).

Atualmente, a tecnologia do cristal líquido (LCD) está dando lugar à tecnologia de LED (diodos emissores de luz). Portanto, é bastante comum encontrar monitores de computador com tecnologia LED ou OLED (LED Orgânico, uma variação da original).

Monitores de LED têm seu funcionamento muito semelhante ao dos monitores de LCD.

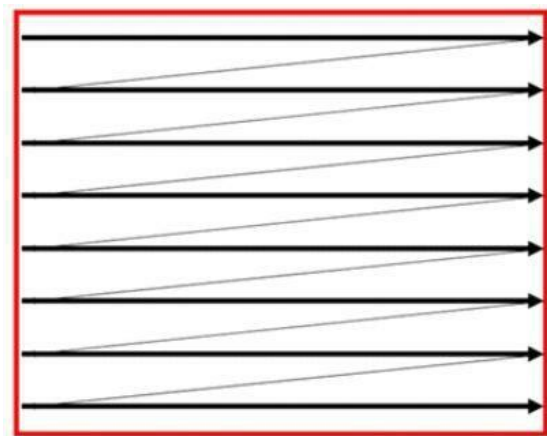
### **Frequência Horizontal e Frequência Vertical**

Num filme ou no jogo, é fácil notar que a imagem (o conteúdo da tela) é modificada o tempo todo (imagens trocam rapidamente nesses dois casos). Isso nos faz concluir que o monitor tem de ficar o tempo todo “redesenhando” a imagem (sobrepondo uma imagem à anterior).

Isso é necessário porque quando a imagem muda, as cores que os pixels apresentam têm de mudar, ou seja, um pixel, que agora está em verde-limão, por exemplo, pode passar a apresentar vermelho-sangue ou azul-coxinha em outro momento e, para isso, ele precisa ser realimentado com a nova cor. (Ahhh... Você nunca viu azul-coxinha? É uma pena!)

Mesmo num trabalho excessivamente estático (como uma janela do Word o tempo todo fixa, com o mesmo conteúdo durante minutos), a tela é redesenhada o tempo todo.

Para “excitar” todos os pontos da tela constantemente, fazendo-os mostrar novas cores ou simplesmente “reacendendo” as cores já existentes, o monitor utiliza dois movimentos: um que faz os feixes de raios catódicos se moverem da esquerda para a direita repetidas vezes e outro que os faz se moverem de cima a baixo, também repetidas vezes. Ao movimento em zigue-zague descendo damos o nome de varredura (é justamente o movimento que é analisado para descrever se um monitor é entrelaçado ou não entrelaçado).



*Figura 2.59 – Exemplo da varredura (desenho da tela).*

À frequência (vezes por segundo) que se pode medir no movimento do feixe da esquerda para a direita damos o nome de **frequência horizontal**. Essa frequência, atualmente, é de cerca de algumas mil vezes por segundo (50 KHz, 60 KHz, 80 KHz). Ou seja, um monitor atual é construído para conseguir fazer o desenho das linhas da tela mais de 50 mil vezes por segundo.



Já a ***frequência vertical*** mede quantas vezes por segundo o feixe desenha a tela de cima até a base. Quando o feixe atravessa a tela de cima para baixo ele simplesmente a desenhou por completo (indo da esquerda para a direita perto de mil vezes). Então, a frequência vertical é muito menor (da ordem de mil vezes menor) que a frequência horizontal – hoje é comum configurar nossos monitores para usarem uma frequência vertical da ordem de 60, 75 ou até 85 Hz (ciclos por segundo).

A frequência vertical também é chamada de ***taxa de atualização da tela***.

## **Resolução**

Uma das coisas que mais se vê nas provas em relação aos monitores é o conceito de resolução. Resolução, caro leitor, ***não é*** qualidade de imagem, fique ciente disso!

Resolução é um conceito meramente numérico – é a contagem dos pixels que estão sendo apresentados naquele momento no monitor.

“Pixel, João?”

Sim! Pixel (abreviação de Picture Element – Elemento da Imagem) é o nome que damos aos pequenos quadradinhos que formam a imagem que a gente vê na tela. Todas as imagens digitais (fotos digitais, por exemplo) são formadas por pixels. São como pequenos “azulejos” na tela. A própria imagem apresentada pelo computador quando se usa o Windows é formada por pixels.



*Figura 2.60* – Até os ícones que vemos (e o ponteiro do mouse) são formados por pixels.

Não costumamos dizer “a tela está apresentando uma imagem com 480.000 pixels”. Em vez disso, o jeito comum de se referir à resolução pela quantidade de pixels na largura x quantidade de pixels na altura. Como em 800 x 600.

Uma resolução de 800 x 600 (também conhecida como SVGA) indica que a imagem da tela está apresentando 800 pixels lateralmente dispostos e 600 pixels verticalmente dispostos. Essa foi, durante muito tempo, a resolução mais comum!

As resoluções mais facilmente encontradas em computador são (ou foram):

- **VGA:** 640 x 480 (a resolução mais baixa) – não é mais utilizada;
- **SVGA:** 800 x 600 (para os padrões atuais, já está sendo considerada baixa) – era a mais comum até bem pouco tempo atrás;
- **XGA:** 1.024 x 768 (pesquisas recentes apontam essa como a mais comum atualmente);
- **SXGA:** 1.280 x 1.024 (normal nos monitores de LCD de 17 polegadas);
- **UXGA:** 1.600 x 1.200 (ainda bastante incomum, mas existe – eu uso essa!).

Todas essas resoluções têm proporção de 4:3 (4 de largura por 3 de altura), perceba isso em 800 x 600, por exemplo. Essa é a proporção para as telas que *não são widescreen* (largura grande, ou “tela de cinema”).

A razão da resolução Widescreen é de 16:9 (16 por 9) ou, em alguns casos, 16:10.

Os valores de **1.280 x 720 (HD)** e **1.920 x 1.080 (Full HD)** são as mais comuns resoluções de alta definição widescreen.

A quantidade de pixels que vemos na tela é gerada pela placa de vídeo, não pelo monitor! Mas o monitor atua como um limitador dessa imagem, determinando qual será a resolução máxima que ele suporta (isso depende de monitor para monitor). Portanto, você poderá usar resoluções muito altas se sua placa de vídeo conseguir desenhá-las e se seu monitor conseguir suportá-las.

“Que tipo de pergunta cai na prova sobre isso, João?”

O mais comum é aquele tipo de questão que exige que você saiba o que acontece com a imagem quando se aumenta ou diminui a resolução. Vamos a um esquema muito simples que o fará entender (causa-efeito) o que acontece quando se altera a resolução da tela.

- Premissa: **Resolução** significa **quantidade de pixels** na tela.

Com base nisso, podemos concluir que:

- Se aumentarmos (↑) a resolução, estamos aumentando (↑) a quantidade de pixels na tela.

Mas aí encontramos um problema, causado pelo monitor (o nosso “castrador”, ou seja, o nosso “limitador”). Em algum momento, antes, durante ou após a alteração da resolução da tela, o seu monitor muda de tamanho?

“Claro que não, João!”

Perfeito, caro leitor! Se o seu monitor permanece com o mesmo tamanho, como é possível colocar mais pixels (aumentar a resolução) se o espaço físico que contém os pixels é exatamente o mesmo?

“Ahh, João... Será que os pixels ficam menores?”

Sim, perfeitamente! Para que a nova (e maior) quantidade de pixels caiba no mesmo espaço físico (o monitor), os pixels têm de diminuir de tamanho (↓). E é isso o que acontece.

Aí encontramos outro efeito: os objetos que são apresentados na tela (ícones, letras, botões, janelas, menus, ponteiro do mouse etc.) são feitos por quantidades fixas de pixels; logo:

- Se o tamanho dos pixels diminui (↓), os tamanhos dos objetos da tela também diminuem (↓). Ou seja, ícones, janelas e tudo mais ficarão menores numa tela que apresenta resolução maior.

“E é para isso que se aumenta a resolução, João? Para os objetos ficarem menores? Não vejo vantagem nisso. Prefiro resoluções menores, então. Prefiro ser capaz de enxergar os objetos!”

Caro leitor, há ainda um último ponto a ser analisado. E esse é justamente a motivação para o aumento da resolução:

- Se os objetos da tela (ícones, janelas) diminuem de tamanho (↓), a área útil da tela (área de trabalho propriamente dita) ficará mais ampla, ou seja, aumentará (↑).

O meu objetivo, ao aumentar a resolução, é conseguir uma tela mais ampla, que me permita colocar mais janelas abertas ao mesmo tempo. A que custo? Ao custo de os objetos da tela ficarem muito pequenos (o que, num monitor maior, como 19 polegadas, não é muito prejudicial).

Resumo sobre resolução (sequência de causa-efeito):

***Resolução = Quantidade de pixels***

↑ ***Resolução***

↑ ***Quantidade de pixels***

↓ ***Tamanho dos pixels***

↓ ***Tamanho dos objetos da tela (ícones, janelas, menus, letras etc.)***

↑ ***Área útil da tela (área de trabalho)***

Se uma dessas setinhas se inverter, todas as outras também irão se inverter!

#### **2.4.7.3. Mouse (entrada)**

É o equipamento que movimenta o ponteiro na tela. Ao mover o mouse por uma superfície plana, seus sensores (que podem ser mecânicos ou ópticos) enviam sinais elétricos desse movimento, e o computador os traduz em movimentos da setinha na tela.



*Figura 2.61 – O mouse óptico.*

Os mouses mecânicos usam uma pequena esfera (bolinha) em sua base para captar os movimentos do equipamento (não são mais tão comuns). Os mouses ópticos usam um sensor luminoso para captar esses movimentos, o que permite muito mais precisão.

Há outros tipos de “dispositivos apontadores” além do mouse. A saber:

- **Touch Pad:** usado em notebooks normalmente, é uma superfície sensível que registra o toque do usuário para servir de indicativo do movimento a ser realizado.



*Figura 2.62 – O touchpad é normalmente usado em notebooks.*

- **Track Ball:** é uma espécie de “mouse de cabeça para baixo”. Nesse dispositivo, o usuário movimenta a esfera e o dispositivo fica parado em relação à superfície.



*Figura 2.63 – Trackball (não se mexe em relação à mesa; a gente gira a bolinha).*

Aproveitando o “tópico” sobre o mouse, eu gostaria de apresentar outro equipamento que não é bem um mouse, mas funciona substituindo-o, especialmente para os profissionais que usam o computador para fazer desenhos (desenhar com o mouse ninguém merece!!!).

- **Tablet (mesa digitalizadora):** é um equipamento em que se usa, normalmente, uma caneta especial para “escrever” sobre uma superfície sensível ao toque. Essa superfície pode, ou não, exibir imagens (ou seja, atuar como monitor, também).



*Figura 2.64 – Tablet Bamboo® da Wacom® (Somente Entrada).*

O exemplo de tablet acima (ou seja, que não exibe imagens) se encaixa no conceito de *dispositivo de entrada*, exatamente como um mouse (na verdade, ela age exatamente como um mouse, pois, ao mover a caneta sobre a superfície, a “setinha” na tela se move).

Por sua vez, os tablets que atuam como monitores (ou seja, exibindo imagens em sua tela sensível), encontram-se na classificação de Dispositivos de Entrada e Saída (ou híbridos) e não só de entrada, como as anteriores.



*Figura 2.65 – Tablet Cintiq® da Wacom® (Entrada e Saída).*

Na prova, por favor, tenha discernimento para interpretar se o elaborador está falando do tablet periférico (esse que estamos vendo agora) ou do tablet “computador”, que nós vimos no início deste capítulo!

#### **2.4.7.4. Impressora (saída)**

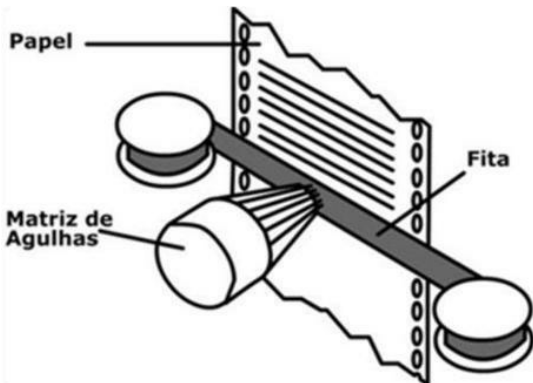
O equipamento que permite que nossos trabalhos sejam postos no papel é a impressora. Há vários tipos e modelos de impressoras atualmente no mercado, mas podemos destacar alguns apenas para fins de estudo.

- **Impressora matricial:** sua técnica de impressão se dá por meio de “agulhas” dispostas em uma matriz. Essas agulhas “batem” numa fita (como na máquina de datilografia), e essa fita, por sua vez, é empurrada contra o papel. Neste tipo de impressão, há contato físico com o papel.





*Figura 2.66 – As impressoras matriciais já foram muito conhecidas do mercado brasileiro (hoje não são muito vendidas).*



*Figura 2.67 – Funcionamento da impressora matricial.*

As impressoras matriciais são pobres em qualidade de impressão, são normalmente muito lentas e extremamente barulhentas. Atualmente, para uso doméstico e corporativo de impressão de documentos são mais utilizados outros tipos de impressoras.

Mas as impressoras matriciais ainda podem ser encontradas em caixas de supermercados (aquelas que imprimem as notas de compra). E, justamente por haver “contato” com o papel (as agulhas batem “di cum força” no papel), elas são usadas quando há necessidade de cópias carbonadas (como notas fiscais e duplicatas de documentos impressos).

- **Impressora jato de tinta:** são as mais comuns hoje em dia. Seu sistema de impressão se baseia em pequenos reservatórios de tinta (cartuchos) que “cospem” a tinta em pontos definidos do papel. A grande maioria das impressoras jato de tinta consegue imprimir em cores.



As impressoras coloridas normalmente possuem dois cartuchos (um preto e um colorido). São reservatórios para as quatro cores primárias: CMYK (ciano, magenta, amarelo e preto). Ciano é o azul-claro, e magenta é rosa. Algumas impressoras apresentam quatro cartuchos separados, um para cada cor.

Lembre-se: há uma diferença entre as chamadas **cores primárias**.

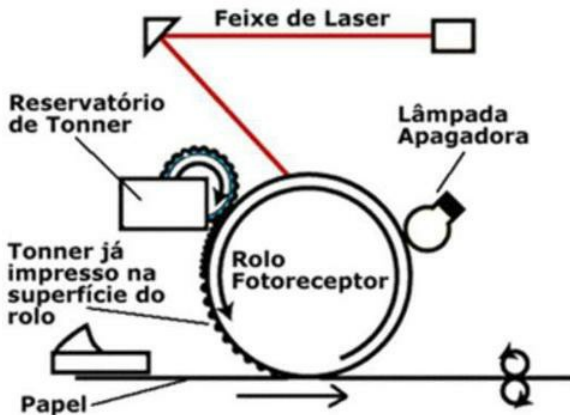
**RGB** (vermelho, verde e azul) são as cores primárias de emissão, usadas pelos **dispositivos de vídeo** (monitor, TV etc.). A partir dessas três cores, qualquer tonalidade pode ser conseguida.

**CMYK** (ciano, magenta, amarelo e preto) é o conjunto de cores primárias de impressão, usadas pelos **dispositivos que imprimem** informações. A mistura correta dessas quatro cores permite conseguir qualquer outra tonalidade numa impressão.

- **Impressora laser:** utiliza um feixe de raio laser para desenhar o objeto a ser impresso em um rolo coberto com um pó chamado tonner. O rolo, por sua vez, se aproxima do papel, e a parte que foi desenhada pelo laser se “prende” no papel devido a uma repulsão por parte do rolo.

Depois de “imprimir” no papel, a superfície do rolo passa por uma lâmpada “apagadora” de modo que seu conteúdo seja limpo para mais um giro em que se repetirá o processo.

Veja o esquema do funcionamento da impressora laser e um exemplo de uma delas a seguir:



*Figura 2.69 – Esquema de funcionamento de uma impressora laser.*



*Figura 2.70 – Exemplo de impressora laser da HP®.*

A maioria das impressoras laser possui apenas uma cor (preto), mas existem impressoras laser coloridas (já estão se tornando mais comuns, apesar de seu preço mais elevado). As impressoras laser monocromáticas (ou seja, imprimem só com tinta preta) já apresentam custo que rivaliza com algumas impressoras jato de tinta, e, mesmo que sejam um pouco mais caras, compensam pelo custo da impressão (o reservatório de tonner consegue imprimir muito mais páginas que um cartucho de tinta).

A principal característica das impressoras jato de tinta e laser é a sua resolução, que é medida em DPI (pontos por polegada). Quanto mais DPI uma impressora tem como resolução, mais qualidade terá o documento impresso. As resoluções mais comuns hoje em dia são 300 DPI e 600 DPI, mas existem alguns modelos de impressoras que conseguem imprimir até 2.800 DPI. (Pelo menos é o que dizem os fabricantes!)

A velocidade de impressão também é uma característica importante: existem impressoras (laser) que conseguem, atualmente, imprimir cerca de 20 PPM (páginas por minuto). Mas a grande maioria não chega nem perto disso (as jato de tinta são muito mais lentas). As

impressoras matriciais são tão mais lentas que sua velocidade é medida em CPS (caracteres por segundo) ou LPS (linhas por segundo).

#### **2.4.7.5. Scanner (entrada)**

Equipamento usado para capturar dados impressos e transformá-los em dados digitais de imagem. Seu uso é muito comum entre profissionais do ramo de design, propaganda, arquitetura etc.



*Figura 2.71 – Scanner comum (scanner de mesa) – também da HP®.*

A principal característica desse equipamento é sua resolução máxima, medida em DPI (pontos por polegada, a mesma medida usada em impressoras). Quanto maior a resolução de um scanner, mais qualidade poderá ter a imagem capturada.

Mas atenção: um scanner possui uma resolução óptica (que é, efetivamente, a resolução do equipamento) e uma resolução intercalada, que é, digamos assim, “melhorada” por software, ou seja, não é real. Na hora de comprar um scanner, analise a resolução óptica (real).

Tudo o que o scanner captura é entendido como imagem, mesmo que a página capturada contenha apenas texto (o que é ruim para esse caso). Já imaginou escanear uma página inteira de um livro, para não ter de digitá-la, e se deparar com a impossibilidade de recortar e copiar

trechos, bem como formatá-los simplesmente porque a página é considerada uma grande e única imagem? Um programa de OCR (Reconhecimento Óptico de Caracteres) resolve esse problema.

Quando compramos um scanner, normalmente ele é acompanhado de alguns programas (edição de foto, melhorias das imagens e OCR). O programa de OCR é capaz de “ler” os caracteres que existem em uma imagem e os transformar em texto editável (colocando-os no Word, por exemplo).

Atenção! Dispositivos leitores de código de barras (também chamados de “scanners” de código de barras) são, também, dispositivos de entrada! Normalmente na forma de “pistolas” (scanners de “mão”).



*Figura 2.72 – Scanner de código de barras.*

Aproveitando: há muitos códigos “de barras” e em outros formatos para registrar informações numéricas e textuais. Os códigos de barras (literalmente, em formato de barras verticais) é usado para registrar, normalmente, números (informações numéricas, como preços, códigos de produtos etc.).



*Figura 2.73 – Código de barras convencional.*

Um código que vem sendo amplamente utilizado hoje em dia e que pode ser visto em vários textos de informática, além de jogos, revistas e até em estacionamentos, é o QR Code (Código QR).

O QR Code é um “código de barras 2D”, que pode ser lido por qualquer dispositivo (celular, smartphone, tablet) que possua uma câmera fotográfica e um programa adequado (note: tem que ter um programa capaz de ler e interpretar o QR Code capturado pela câmera!).

Um QR Code pode conter mensagens de texto bastante complexas (mais de 1.000 caracteres), como um endereço de Internet, uma mensagem “secreta”, entre outros! Para ler o QR Code a seguir, use seu smartphone ou tablet, com o programa próprio para leitura de QR code, e aponte a sua câmera para a imagem seguinte!





*Figura 2.74 – QR Code (Mensagem secreta! Descubra-a!).*

Note que, tanto o QR Code quanto o código de barras (ou qualquer outro código desse tipo) é capaz, sim, de armazenar informação codificada. Mas para lê-la, é necessário um dispositivo (periférico, como um scanner ou uma câmera) e um programa adequado!

#### **2.4.7.6. Multifuncional (entrada e saída)**

Eis um dispositivo muito comum nos últimos anos: o “multifuncional” (para alguns, como a FCC, a “impressora multifuncional”), que é, simplesmente, um scanner “preso” no topo de uma impressora.

Esse equipamento alia as características de captura do scanner com a capacidade de impressão da impressora, criando um grande dispositivo capaz de fazer entrada e saída. Capaz, também, de atuar como uma máquina copiadora.



*Figura 2.75 – Uma multifuncional (scanner + impressora + copiadora).*

Há multifuncionais que imprimem com tecnologia laser, embora as mais comuns sejam, claro, as que imprimem por jato de tinta.

#### **2.4.7.7. Modem (entrada e saída)**

O modem (modulador/demodulador) é um equipamento de comunicação que permite que dois computadores fiquem conectados (troquem informações) através de uma linha de transmissão de sinais analógicos (normalmente a linha telefônica).



*Figura 2.76 – Placa de fax/modem (já não é tão comum!).*

A função do modem é traduzir os pulsos elétricos digitais (existentes no interior do computador) em variações elétricas analógicas (forma de transmissão dos dados na linha telefônica). Quando um modem realiza o processo de tradução digital-analógico, dizemos que ele está realizando uma **modulação**. Quando o modem faz o processo inverso (analógico-digital), essa tradução é chamada de **demodulação**.

Atualmente, os modems convencionais (para linhas telefônicas) atingem uma taxa de transferência de **56 Kbps** (leia-se **56 Kilobits por segundo**). Não há modems telefônicos, nem em projetos futuros, que atinjam valores superiores. A verdade é que esse equipamento está praticamente MORTO. Quanto mais o acesso à Internet se desvincula da linha telefônica, mais o modem telefônico (convencional) se torna desnecessário.

Já existem outros equipamentos de conexão em rede que usam outros sistemas que diferem da linha telefônica (ADSL, cabo, 3G etc.), que serão discutidos posteriormente.

O que importa é que: se é modem, não importando o “tipo” de modem que é, pode ter certeza de que é classificado como periférico de **Entrada e Saída**. E, também, todo modem, não

importando o tipo, faz modulação e demodulação!

#### 2.4.7.8. Placa de rede (entrada e saída)

A comunicação entre computadores não se dá somente através da linha telefônica (ou de outros sistemas de longa distância). É possível ligar vários equipamentos em redes locais, dentro das casas e empresas. Uma rede local, também chamada LAN, exige certos equipamentos específicos, como cabos especiais, hubs, switches (todos discutidos no capítulo sobre Redes).

Além desses, é necessário que cada computador possua um pequeno equipamento capaz de se comunicar através dessa estrutura de cabamentos. Esse equipamento chama-se placa de rede, ou adaptador de rede. Outro nome comum para ele é NIC (Placa de Interface de Rede).



Figura 2.77 – Placa de rede Ethernet (interna ao gabinete).

Uma das arquiteturas de rede (veremos posteriormente o que isso significa) mais usadas hoje em dia é a arquitetura Ethernet. A maioria dos equipamentos (cabos, placas, concentradores etc.) para redes é construído seguindo essa arquitetura de funcionamento. Daí a razão por que, atualmente, as placas de rede são normalmente chamadas de placas Ethernet.

Só lembrando: **Ethernet** é, hoje, “sinônimo” de **LAN com fios**. Isso se dá porque hoje em dia, Ethernet é a tecnologia mais usada, disparadamente, no mundo todo, para fazer redes locais (LAN) com fios (cabeadas)!

Já houve várias gerações (padrões) de Ethernet ao longo da “história”. O padrão mais usado

hoje (ainda, embora experimentando seu “crepúsculo”, ou seja, sua “aposentadoria” próxima) é o **Fast Ethernet**, no qual as placas de rede são construídas para atingir até 100 Mbps (Megabits por segundo) de velocidade de transferência. Uma placa construída em um padrão mais novo consegue se comunicar com placas mais antigas, mesmo se suas velocidades não forem iguais.

Por essa razão, as placas de rede atuais vêm com a seguinte inscrição: Placa de rede 10/100 (ou placa Ethernet 10/100). Isso indica que a placa pode se conectar a 100 Mbps ou a 10 Mbps (que é a velocidade da geração anterior, a Ethernet original), dependendo da necessidade.

A geração sucessora do padrão Fast Ethernet é 10 vezes mais rápida: o **Gigabit Ethernet** consegue 1.000 Mbps (ou 1 Gbps – 1 Gigabit por segundo). Todas as placas de rede atualmente vendidas (em laptops, desktops e afins) é nessa velocidade! Não há mais produtos novos no mercado que ofereçam placas de rede com fio de velocidades anteriores.

Lembre-se: se duas placas de rede de velocidades diferentes estão “conversando”, a velocidade dessa conversa é definida pela placa mais lenta. Isso é lógico porque quem é rápido pode diminuir a velocidade, mas quem é lento não consegue subir.

#### **2.4.7.9. Placa de rede Wi-Fi (entrada e saída)**

Outra placa de rede muito comum nos dias de hoje (especialmente em equipamentos portáteis, como notebooks, ultrabooks e netbooks) é a placa Wi-Fi, que nada mais é que um dispositivo de comunicação que permite ao computador se conectar a uma rede Wi-Fi (Wireless).



*Figura 2.78 – Placa de rede Wi-Fi.*

As placas Wi-Fi também são classificadas segundo seus padrões. Todas as tecnologias usadas para Wi-Fi (redes locais sem fio) são padronizadas segundo as normas do IEEE (Instituto de Engenheiros Elétricos e Eletrônicos, uma instituição que padroniza tecnologias ligadas à informática), sob a norma 802.11.

Ou seja, falou em **802.11**, falou em **Wi-Fi**! Os padrões mais importantes de serem lembrados são os seguintes:

- **802.11b:** o padrão mais antigo. Os equipamentos que trabalham neste padrão usam uma frequência de 2,4 GHz e transmitem dados a 11 Mbps.
- **802.11g:** também utiliza a faixa de frequência dos 2,4 GHz. Transmite dados a 54 Mbps.
- **802.11a:** usa a faixa de frequência de 5 GHz para transmitir a 54 Mbps.
- **802.11n:** é um padrão recente e está fazendo um grande sucesso. Garante transmissões da ordem de 300 Mbps (três vezes mais que o Fast Ethernet), usando as duas faixas de frequência possíveis (2,4 GHz e 5 GHz). Esse é o padrão mais usado e comercializado hoje em dia!

Mas, como era de se esperar, já há uma novíssima geração chegando:

- **802.11ac**: esse padrão traz transmissões a 5 GHz com velocidades de, pasme, 1.300 Mbps (1,3 Gbps). Isso é mais do que a rede Gigabit Ethernet! Esse padrão, porém, é tão novo que não há sequer a homologação (autorização) deste formato e apenas uma fabricante começou a produzir equipamentos para ele.

#### 2.4.7.10. Placa de som (entrada e saída)

É o equipamento capaz de transformar as informações digitais dos programas e jogos em som estéreo (para sair nas caixinhas de som) e para transformar os sons capturados de um microfone (ou de um instrumento musical, CD player etc.) e transformá-los em informações digitais para serem processadas pela CPU do computador.



*Figura 2.79 – Placa de som.*

Há vários modelos de placas de som: desde as mais simples até aquelas mais caras e robustas (como a da imagem acima, que possui até mesmo a saída óptica – conectores quadrados na extremidade – usada em home cinema!).

#### 2.4.7.11. Placa de vídeo (saída)

É o equipamento responsável por “desenhar” os dados que aparecem no monitor do computador. Todos os dados que saem da CPU em direção ao monitor passam pela placa de vídeo, que converte os sinais elétricos digitais em sinais RGB (as cores primárias). Sendo assim, o monitor já recebe os sinais da maneira como deve emití-los para o usuário.



*Figura 2.80 – Placa de vídeo.*

As placas de vídeo atuais trazem, consigo, sua própria “CPU” (chamada, na verdade, de “GPU” – ou Unidade de Processamento Gráfico). A GPU trabalha desenhando a imagem que o usuário vê, determinando a cor de cada pixel da tela.

Algumas placas de vídeo “mais simples” (normalmente nos netbooks e laptops menos caros) não possuem GPU, portanto utilizam o poder de processamento que a CPU do computador dá, tornando-os um pouco menos indicados para assistir a filmes e jogos.

#### **Memória de Vídeo**

Todas as placas de vídeo (normais ou não) possuem memória, a chamada memória de vídeo,



que nada mais é que uma memória DRAM usada somente para armazenar os dados de imagem (pixels e suas cores). Quanto mais memória de vídeo uma placa dessas possuir, mais resolução e mais cores a imagem do computador pode apresentar; portanto, o principal definidor da qualidade da imagem é a placa de vídeo, e não o monitor.

Atualmente são comuns placas de vídeo normais com 256 MB de memória. As placas 3D exigem memórias de 512 MB até 2 GB. Mas, na realidade, se usarmos o computador apenas para tarefas simples, como digitar no Word e navegar na Internet, não precisaríamos de mais que uma placa de vídeo com 16 MB... Não mais que isso!

### As Saídas da Placa de Vídeo

As placas de vídeo atuais podem ser dotadas de algumas saídas (encaixes) específicas para monitores de tipos específicos. Vamos conhecê-las:

- **Saída VGA:** a saída normal (que aparece normalmente em azul), ligada em qualquer monitor analógico (monitores de computador – atuais e antigos – podem ser ligados no conector VGA). Serve tanto para monitores de LCD, LED, quanto para os monitores de CRT. É um conector analógico (os sinais que trafegam por ele são analógicos).
- **Saída DVI:** saída digital (normalmente um conector branco), usada para conectar monitores (e TVs) de LCD e LED mais modernos – oferece uma qualidade superior de imagem para esses dispositivos.
- **Saída S-Video:** uma saída redonda, usada para ligar a placa de vídeo diretamente a uma TV comum (antiga).
- **Saída HDMI:** uma saída digital, normalmente usada para conectar o computador diretamente em TVs LCD e LED da nova geração. Oferece qualidade de imagem tão boa quanto a DVI (ou até melhor!).

O conector HDMI normalmente transporta vídeo e SOM, mas a maioria das placas de vídeo só o utiliza para transmitir sinais de imagem (vídeo).

- **Saída Display Port:** usada em apenas alguns tipos de placas de vídeo e computadores (é padrão nos computadores e laptops da Apple), este sistema digital de vídeo é promessa para o futuro.



*Figura 2.81 – Saídas comuns na placa de vídeo (HDMI, DVI e VGA).*



*Figura 2.82 – Conector Mini Display Port, em computador Apple®.*

Com isso, meus amigos, terminamos a análise dos mais importantes dispositivos de entrada e saída (periféricos) do computador. Vamos partir para o próximo “tópico” do nosso estudo, que são os barramentos presentes no computador.

#### **2.4.8. Barramentos**

Os barramentos do computador são divididos em basicamente dois grandes grupos: o barramento de sistema (atualmente substituído pelo chipset da placa-mãe) e os barramentos de expansão.

##### **2.4.8.1. Barramento de sistema**

Vamos começar estudando o barramento de sistema, embora este pareça não existir!

“Mas, João, foi você mesmo quem disse que ele não existe!”

Ele não tem um “corpo” exatamente do jeito como é desenhado, caro leitor. Ou seja, não é um “barramento” no sentido visual da coisa. Ele não se parece com um barramento (ou seja, não se parece com uma “estrada”).

O barramento de sistema, mesmo não sendo mais visível assim, é dividido em três sub-barramentos (ou três conjuntos de linhas) conhecidos como barramento de dados (ou linhas de dados), barramento de controle (ou linhas de controle) e barramento de endereços (ou linhas de endereços).

A imagem que representa essa distinção é mostrada a seguir:

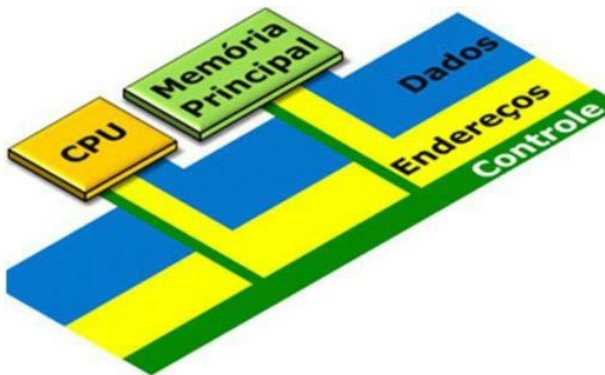


Figura 2.83 – Barramento do sistema.

Cada conjunto de linhas é o caminho por onde trafega um tipo específico de informação.

“É, já sei, e preciso saber o que cada um faz, não é?”

Exatamente, caro leitor! Especialmente para a ESAF e a FGV (Fundação Getúlio Vargas), que exigem esse tipo de conhecimento; podemos encontrar várias perguntas sobre isso em outras bancas, mas normalmente quando o concurso é “mais pesado”, como para Auditores Fiscais estaduais, entre outros. Vamos conhecer o que cada um desses barramentos é capaz de fazer:

- **Barramento de dados:** essa parte do barramento de sistema é responsável por transferir **dados e instruções** pertencentes aos programas que estão sendo executados no computador naquele instante.

A largura do barramento de dados (32 ou 64 bits) determina a palavra daquele processador. Ou seja, um processador só é considerado de 64 bits se seu barramento de dados possuir essa largura. Hoje em dia, praticamente todos os processadores possuem barramento de dados de 64 bits.

Nesse barramento, são transferidas informações nos dois sentidos (CPU → Memória e Memória → CPU), ou seja, ele é bidirecional.

- **Barramento de endereços:** por ele são transferidos os **endereços** das posições na memória principal que serão acessadas naquele momento.

A largura do barramento de endereços determina a capacidade de endereçamento (gerenciamento de memória) de um processador. Ou seja, quanto mais largo for o barramento de endereços de um computador, **mais memória principal ele pode ter**, pois mais posições de memória ele poderá endereçar.

A capacidade de endereçamento de memória diz respeito à quantidade de posições de memória que uma CPU é capaz de gerenciar na memória RAM. Para obter o número exato de posições de memória possíveis usa-se a expressão:  $P = 2^K$ . Onde P é o número de posições (células) de memória principal e K é a largura do barramento de endereços (em bits).

Portanto, para um barramento de endereços de 32 bits (a maioria atualmente), a quantidade de células de memória que se pode ter é  $2^{32}$ , ou seja, cerca de 4 bilhões de células! E como cada célula tem capacidade para apenas 1 byte, é correto afirmar que a maioria de nossos computadores só pode gerenciar cerca de 4 bilhões de bytes (4 gigabytes).

Essa é a razão de nossos micros não poderem ter mais de 4 GB de memória RAM: a largura do barramento de endereços. Claro que já há processadores com barramentos de memória maiores (40 bits, por exemplo, que dão até 1 TB de capacidade de memória máxima), mas os sistemas operacionais (como o Windows, por exemplo) também são limitantes para o tamanho máximo da memória RAM!

Esse barramento permite a comunicação apenas no sentido CPU  $\square$  Memória, porque a CPU é quem determina qual endereço vai acessar na memória e o fará enviando tais sinais por esse barramento. A memória nunca enviará sinais de endereçamento (a memória nunca “envia” ou “ordena” nada!).

- Barramento de controle: por ele são transferidos os sinais de controle que a CPU envia para os demais componentes do micro ou recebe deles.

A largura do barramento de controle é simplesmente desprezível, pois não determina nenhuma característica útil ao computador.

Ou seja, se a CPU fosse o “Coronel Jesuíno” (personagem da obra literária *Gabriela*), ela usaria o barramento de controle para dizer “Memória... se prepare que eu vou lhe usar!”. Da mesma forma, pelo barramento de controle, outros dispositivos do micro, como os periféricos de entrada, interrompem a CPU para dizer “Ei, CPU, pare o que está fazendo e preste atenção em mim!” (essas “interrupções” são chamadas de... de... Interrupções!).

Em suma, é pelo barramento de controle que os periféricos enviam os sinais de interrupção para a CPU!

E... Falando em periféricos... Vamos agora estudar os barramentos “secundários” em importância para o micro: os barramentos de expansão, que servem para ligar os periféricos ao micro.

Lembre-se somente de que o barramento de sistema está na placa-mãe, sob a forma de um chip (a ponte norte).

Esse barramento também é bidirecional. A CPU pode tanto enviar dados de controle por ele quanto pode receber (mas não da memória!). A CPU consegue receber sinais de controle vindos de outros dispositivos, como os periféricos de entrada.

A memória apenas recebe sinais de controle, ela não envia sinais a nenhum outro dispositivo. Ou seja, a memória NÃO MANDA... Ela só obedece! (Manda quem pode, obedece quem tem juízo, ou... Quem não tem escolha!).

#### 2.4.8.2. Barramentos de expansão

Barramentos de expansão são, como visto anteriormente, as vias que fazem a informação trafegar entre o chipset e os periféricos do computador. Dentro dessa classificação podemos citar duas subdivisões: os barramentos internos e os barramentos externos.

Os barramentos internos são aqueles que ligam o chipset aos equipamentos localizados dentro do gabinete. Existem vários tipos de barramentos para ligar os equipamentos internos, até porque existem vários equipamentos internos diferentes (e, mais importante, com formatos diferentes), como modems, placas de vídeo, HDs, gravadores de DVD etc.

Os barramentos de expansão também fazem parte da placa-mãe do computador e apresentam-se na forma de seus conectores (os slots e as portas visíveis a olho nu). Ou seja, não dá para ver realmente os barramentos, mas apenas seus slots e portas (os conectores nos quais encaixamos os cabos e dispositivos periféricos).

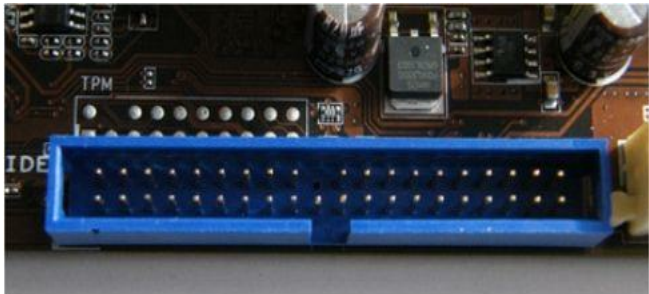
Vamos a eles:

##### **Barramento IDE**

ERA usado para conectar as unidades de armazenamento internas (HD, drive de CD, gravadores de CD, drives de DVD etc.) à placa-mãe do computador. O barramento IDE tinha largura de 32 bits.

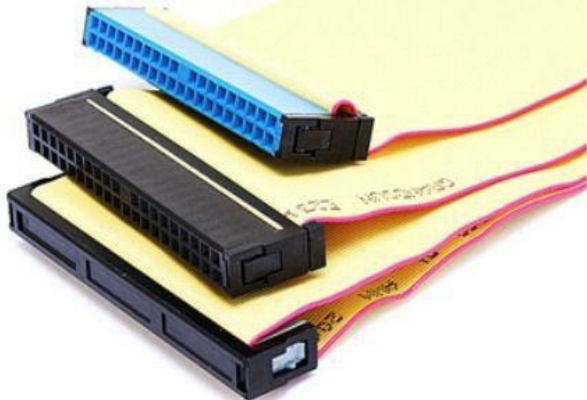
Cada barramento IDE permitia a conexão de apenas dois equipamentos de disco. Mas como já foi comum haver dois barramentos IDE (chamados de IDE primário e IDE secundário) em um computador o total de equipamentos de armazenamento interno chegava a quatro.

Os parágrafos acima descrevem o barramento IDE sempre no passado... Pois é... Ele não existe mais! Nas placas-mães atuais, não há mais barramento IDE (também chamado de **PATA**, ou **ATA Paralelo**) – ele foi totalmente substituído pelo ATA Serial (Serial ATA), ou SATA, que veremos a seguir!



*Figura 2.84 – Slot IDE.*

Os equipamentos ligados aos barramentos IDE eram conectados a esse através de um cabo denominado ***Cabo Flat*** (Cabo “achatado”).



*Figura 2.85 – Cabo Flat.*

### *Mestre e Escravo*

Existe uma nomenclatura que está longe de ser politicamente correta em relação ao barramento IDE: **mestre** e **escravo**. O que isso significa? Quando havia dois dispositivos (HDs, CDs, DVDs etc.) ligados em um mesmo barramento IDE, um deles era denominado mestre (master) e o outro recebia a “carinhosa” denominação de escravo (slave).

Normalmente, se definia quem era quem em cada dispositivo, ou seja, um HD era montado em um micro já sabendo se seria mestre ou escravo. Isso é possível realizando a configuração do dispositivo em seu painel traseiro através de jumpers (pequenos conectores ligados a pinos metálicos, fechando e abrindo circuitos). Ou seja, as configurações de mestre e escravo eram manuais (nada automáticas!).

Como havia dois canais (barramentos) IDE independentes em um micro, podemos concluir que os discos IDE instalados em um computador podiam ser configurados como: primary master (mestre primário) e primary slave (escravo primário) – ambos ligados no canal IDE primário – e secondary master (mestre secundário) e secondary slave (escravo secundário) – ligados ao barramento IDE 1, chamado de canal secundário, ou seja, o segundo barramento IDE.

Então, quando alguém disser “Instalei um disco de 80 GB como mestre primário e deixei o CD-ROM como mestre secundário” não será nenhum mistério, será? Mas, claro, ninguém mais



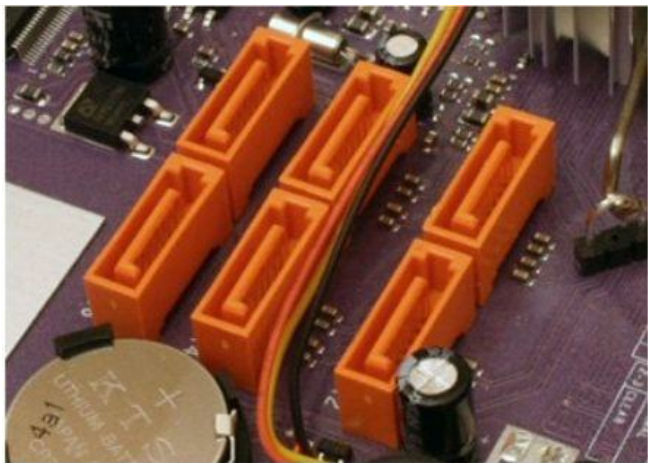
dirá isso, porque o IDE já era!

### **Barramento SATA (Serial ATA)**

Eis o alçoz do IDE! O barramento SATA é a razão de as placas-mãe, hoje em dia, não trazerem mais slots IDE! (Na verdade, já há alguns anos!)

SATA é um barramento serial (isso quer dizer que tem largura de 1 bit) que traz inúmeras vantagens em relação ao IDE. A primeira delas, a velocidade: um barramento SATA original (em 2002, aproximadamente) possuía velocidade de transferência da ordem de 150 MB/s, mas hoje, no SATA II, já se pode instalar um HD com velocidade de transferência de cerca de 300 MB/s. (Essa “geração” é normalmente conhecida como SATA 3 Gbps, embora efetivamente só transfira 2,4 Gbps – 300 MB/s). É a mais comum da atualidade!

Hoje, porém, também já é possível encontrar (embora mais raramente) discos rígidos, gravadores de DVD e Blu-ray e SSDs fabricados com o barramento SATA III, ou SATA 6 Gbps, que transfere dados a 600 MB/s. Esta velocidade, porém, só poderá ser inteiramente aproveitada se a placa-mãe, óbvio, possuir barramento SATA nesta geração!



*Figura 2.86 – Seis conectores SATA numa placa-mãe.*

Num barramento SATA não dá para ligar mais de um disco (como no IDE), mas, em

compensação, é possível haver, numa placa-mãe, diversos barramentos SATA (já que são muito simples de fabricar, pois só usam, tecnicamente, um único fio). Há placas-mãe com seis, oito e até mesmo doze conectores SATA disponíveis.

Os dispositivos SATA admitem a técnica de hot swap (“troca a quente”), que, essencialmente, significa que os discos nessa tecnologia podem ser conectados e desconectados do computador com a máquina ligada, sem risco de dano para o micro ou para o disco (eu ainda não tentei isso, é verdade, porque não tive coragem! Mas dizem que funciona!!).

### *eSATA*

Uma variação do SATA usada em discos rígidos externos é chamada de eSATA (External SATA – SATA Externo).

Na verdade, não se trata de um barramento diferente, mas de uma “extensão” do SATA. É um fio, conectado a algum slot SATA interno que fornece “portas” externas para a conexão de HDs e gravadores de DVD externos.

Hoje em dia é muito comum instalar HDs em “cases” (caixas) para torná-los externos (removíveis), e uma das formas de conectá-los ao computador é por meio do slot eSATA.



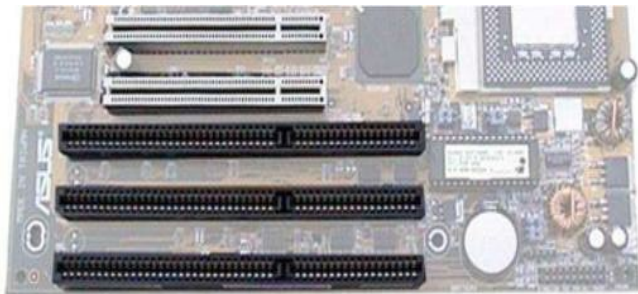
*Figura 2.87 – Conector eSATA (o de baixo) localizado na lateral de um laptop.*

Uma coisinha a mais: hoje em dia, é normal, pelo menos nos laptops, que a porta eSATA também seja utilizável por dispositivos USB, ou seja, a porta (a da figura acima é um exemplo)

é, na verdade, um “combo” (combinação) de USB com eSATA. É uma porta só, mas dá para encaixar equipamentos eSATA e equipamentos USB nela (porque os dois formatos são “parecidos”).

### **Barramento ISA**

Muito comum em micros mais antigos para encaixar placas de expansão, como modems, placas de som, placas de vídeo, entre outros, o barramento ISA já não é mais fabricado em placas-mãe desde 1998 mais ou menos, por ser muito lento em relação às novas tecnologias.



*Figura 2.88 – Três Slots ISA (esta placa é de 1997).*

O barramento ISA não possui a característica plug and play, ou seja, qualquer equipamento conectado a esse barramento deve ser instalado no computador através de um processo manual (e muitas vezes traumático). O Windows não reconhece automaticamente a presença de equipamentos conectados a esse barramento, como faz em outros barramentos do computador.

### **Barramento PCI**

O PCI chegou para ser o substituto do barramento ISA, mas durante muito tempo teve de conviver com este nas placas-mãe do mercado. O barramento PCI também é usado para qualquer tipo de equipamento em formato de placa (placas de expansão), como modem, placa de rede, placa de som, placa de vídeo e afins.



*Figura 2.89 – Dois Slots PCI (os brancos).*

Atualmente, não é possível encontrar nenhum ISA nas placas-mãe, mas é possível, ainda, encontrar alguns PCI (a foto acima é de uma placa-mãe com, no máximo, 2 anos de fabricação). O barramento PCI é muito mais veloz que o barramento ISA, e esta é uma das principais razões de tê-lo substituído.

O barramento PCI é plug and play, característica que permite que os componentes encaixados a esse barramento sejam automaticamente detectados pelo sistema operacional do computador (Windows, no nosso caso). Ou seja, instalar um equipamento qualquer nesse barramento é muito menos trabalhoso que no barramento ISA.

### **Barramento AGP**

Foi um barramento criado *apenas para uso por placas de vídeo*. Sua taxa de transferência era muito superior à do barramento PCI (que também era usado anteriormente para as placas de vídeo). Lembre-se: não se encontram outros equipamentos (como modem, placa de rede, placa de som) ligados ao AGP; esse barramento foi feito para conectar apenas placas de vídeo.

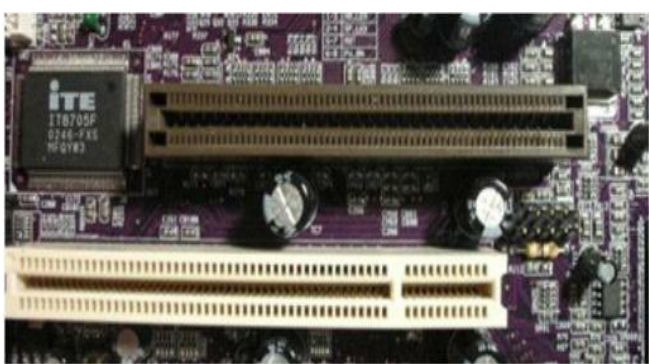


Figura 2.90 – Slot AGP (o marrom, em cima).

Foi um barramento que esteve presente em quase todas as placas-mãe, com exceção daquelas que já traziam a placa de vídeo on-board, nas quais ele foi dispensado com a desculpa de não ser necessário.

“Ei, João, você usou o parágrafo no tempo passado!”

Precisamente, caro leitor! ***O AGP também já era!*** Apesar de ainda existirem algumas placas-mãe (as antigas) com esse barramento, é mais comum encontrar, hoje em dia, o seu sucessor nas placas novas. Vamos conhecê-lo mais à frente.

### **Barramento PCI Express**

Eis o culpado pela aposentadoria do PCI e do AGP! O PCI Express é o barramento atualmente mais “em moda” nas placas-mãe mais modernas (desde 2006).

O barramento PCI Express é serial (sim, serial!) que pode ser usado para conectar qualquer tipo de equipamento em forma de placa (modem, placa de rede, placa de vídeo, placa de som etc.).

Embora seja uma conexão ponto a ponto (cada dispositivo está ligado ao seu próprio caminho, sem “compartilhá-lo” com nenhum outro dispositivo), o que, em si, contraria a ideia de barramento, essa tecnologia é comumente chamada de barramento PCI Express.

Um slot PCIe (abreviação do PCI Express) pode ser montado em várias configurações, de acordo com o número de linhas seriais conectadas a cada slot (1, 4, 8 ou 16 linhas). Essas linhas são reunidas em grupos e conectadas aos slots na placa-mãe, gerando, assim, as diversas variantes de PCIe (do PCIe x1 ao x16).



*Figura 2.91 – Três Slots PCI Express x1 (os brancos pequenos).*

Cada linha PCIe consegue transmitir dados a uma velocidade de cerca de 250 MB/s. Em um slot com 16 linhas (chamado de PCIe x16), atinge-se, claro, 16 vezes mais velocidade (cerca de 4 GB/s).

Apesar de haver quatro possíveis configurações (x1, x4, x8 e x16), é muito comum encontrar, nas placas-mãe atuais, apenas o x1 e o x16. Sendo que o x1 é usado para qualquer tipo de placa de expansão (portanto, atua como o substituto do PCI) e o x16 é usado somente para placas de vídeo (já que elas exigem mais velocidade), assumindo, assim, o “status” de substituto do AGP.

E o AGP já morreu mesmo! Em quase a totalidade das placas-mãe atuais (as que são vendidas no mercado hoje em dia), usa-se PCIe x16 em vez de AGP.

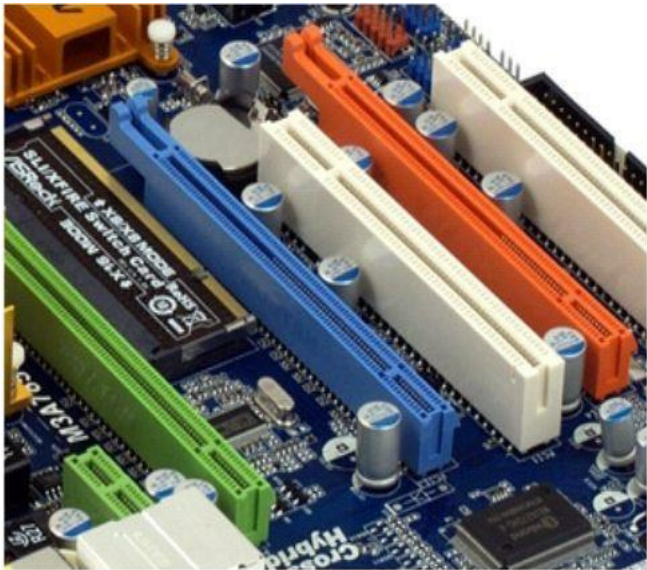


Figura 2.92 – Três Slots PCIe x16 (os três mais compridos – coloridos) – os brancos são PCI (antigos).

Então não se esqueça disto, caro leitor: PCI Express (PCIe) é um barramento serial. Cada placa-mãe vem com várias linhas (em média 32 linhas) que podem ser combinadas de diversas formas em vários slots (isso depende do modelo de placa-mãe).

Um exemplo de combinações possíveis: numa placa-mãe que tenha circuito controlador de 48 linhas PCIe, os slots poderiam estar distribuídos desta forma: 2 PCIe x16, 1 PCIe x8, 1 PCIe x4 e 4 PCIe x1 (totalizando 48 linhas).

E apesar de ser possível realizar esses vários tipos de junções de linhas na placa-mãe, as duas configurações mais comuns de PCI Express são os extremos (x1 – que usa apenas uma linha – e x16, que usa 16 linhas e é usado *apenas para placas de vídeo*).

## Barramento SCSI

Este barramento definitivamente não é comum entre os computadores pessoais dos usuários domésticos. O barramento SCSI (lê-se “iscâsi”) é muito versátil, podendo ser encontrado para conectar diversos tipos de equipamentos, como scanners e impressoras, por exemplo, mas é um exímio concorrente para o barramento SATA. O SCSI é muito usado em servidores de empresas, que normalmente precisam de uma maior velocidade de conexão com os discos rígidos, CDs, unidades de fita etc.

Uma das vantagens do barramento SCSI em relação ao IDE e ao SATA é o fato de poder conectar até 15 equipamentos ao mesmo tempo (contra apenas dois do IDE), o que permite uma expansão da capacidade de armazenamento do computador muito mais facilitada. Outra vantagem do barramento SCSI é a taxa de transferência, que pode chegar até a 320 MB/s (hoje, porém, o SATA já consegue atingir velocidades muito superiores).

Não é comum encontrar slots SCSI em placas-mãe; portanto, utiliza-se uma placa de expansão para que esta consiga conectar-se aos equipamentos SCSI. Essa placa é chamada **controladora SCSI**.



*Figura 2.93 – Placa controladora SCSI.*

O barramento SCSI é originalmente paralelo (possui largura de 16 bits). Mas, atualmente, existem mais vantagens em barramentos seriais que nos paralelos. Um barramento serial pode ser muito mais veloz que um barramento paralelo, pois nos paralelos há limitações impostas pela



indução eletromagnética que um fio faz nos seus vizinhos (há muito ruído que atrapalha a transmissão).

Os servidores atuais estão utilizando uma “variação” serial do SCSI: o **barramento SAS** (Serial Attached SCSI – SCSI Anexado Serial). SAS é uma conexão ponto a ponto (ou seja, não exatamente um barramento) que permite a ligação de discos SCSI em um ambiente de conexão serial.

Há diversas vantagens no SAS em relação ao SCSI tradicionalmente paralelo, mas não as discutiremos aqui, pois tanto o SCSI (hoje considerado antigo) quanto o SAS são para servidores (computadores centrais das empresas), pelo custo e pela velocidade que oferecem.

Para você memorizar: digamos que o SAS está para o SCSI assim como o SATA está para o IDE (PATA). São evoluções seriais de antigos e famosos barramentos paralelos usados para discos.

## **Barramento PS/2**

É o barramento atualmente utilizado para conectar mouse e teclado. É um barramento lento (transfere dados com pouca velocidade) e funciona de forma serial (ou seja, sua largura é de 1 bit apenas – ele transfere um bit por vez). Há duas portas na parte traseira do gabinete, uma para o mouse e a outra para o teclado.



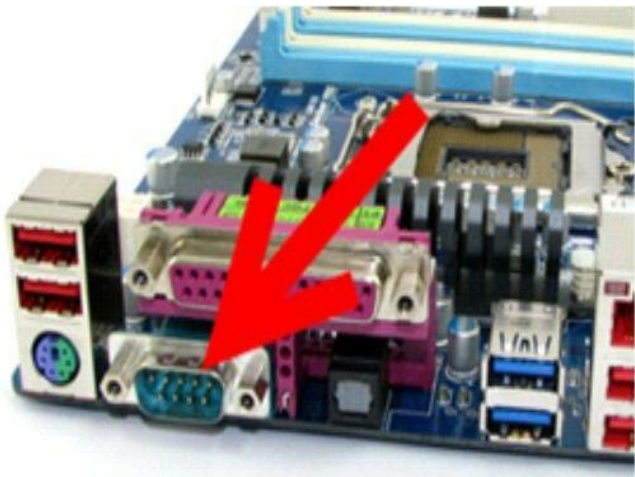
*Figura 2.94 – Porta PS/2 (apenas uma).*

## **Barramento Serial**

É um barramento usado por uma série de equipamentos que transferem relativamente pouca informação, como mouses, modems, câmeras (webcam) etc. Quando não havia porta PS/2 no micro, o mouse normalmente era conectado a uma das portas seriais disponíveis.

O barramento serial foi padronizado seguindo-se a norma técnica RS-232 (padrão oficial

projetado pela EIA – Electronic Industries Association). Ou seja, qualquer prova que citar o RS-232 está se referindo à porta serial comum de um computador.



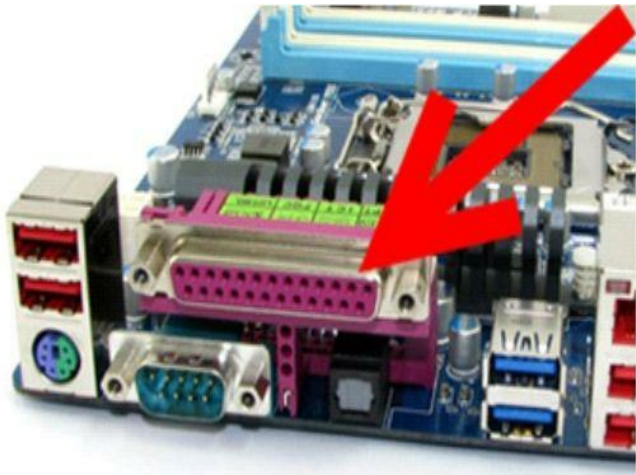
*Figura 2.95 – Porta serial (conector DB-9).*

Num computador pessoal, os conectores mais comuns para o barramento serial são o DB-9 (que usa nove pinos, apontado na figura anterior), o DB-15 (que usa 15 pinos) e, o mais antigo, o DB-25 (25 pinos). O barramento serial utiliza apenas um único canal de transmissão de dados (largura de 1 bit). Ou seja, os bits são transmitidos em fila, um a um, daí o nome barramento SERIAL (em série).

### **Barramento Paralelo**

É um barramento bastante antigo, que, como o serial, basicamente não há mais em computadores atuais. A porta paralela usa conector DB-25 (antigamente usado pelo barramento serial).

A porta paralela (um computador tinha normalmente uma) era usada para conectar equipamentos que exigiam um tráfego de dados mais intenso, como impressoras, scanners, unidades de armazenamento externas (como os antigos zip drive, por exemplo) etc.



*Figura 2.96 – Porta paralela (conector DB-25).*

### **Barramento USB**

Sem dúvida alguma, é o barramento externo mais utilizado atualmente. Atualmente todo tipo de equipamento periférico externo é ligado pela porta USB. O barramento USB (Universal Serial Bus – Barramento Serial Universal) é o substituto dos barramentos serial, paralelo e PS/2.

Um computador atual pode conter diversas portas USB em sua traseira (normalmente de 4 a 10), o que permite a conexão de diversos equipamentos, como impressoras, scanners, teclados.



*Figura 2.97 – Portas USB (quatro delas).*

Uma característica muito interessante sobre o barramento USB é que a ele podem ser conectados 127 equipamentos diferentes em fila, ou seja, um ligado ao outro. Já imaginou? Seu micro conecta-se à impressora, que se conecta ao monitor, que se conecta ao scanner, que se conecta ao teclado, que se conecta ao... E por aí vai!

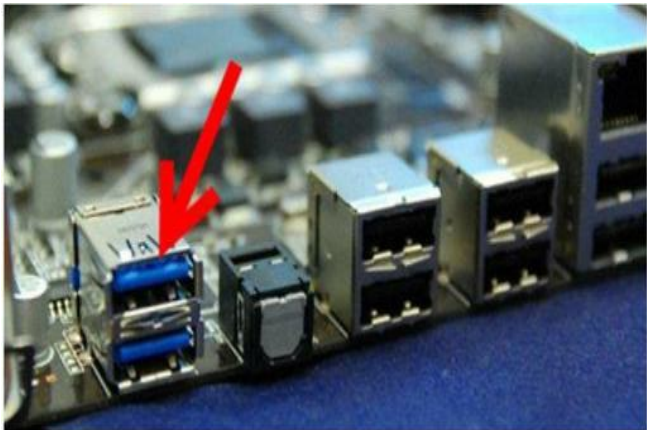
Não precisa ser exatamente assim; você pode conectar ao seu computador um equipamento que vai funcionar como um “T” (desses benjamins de tomada elétrica mesmo). Esse equipamento, chamado **Hub USB**, tem a finalidade de se conectar a uma porta e fornecer várias portas para outros equipamentos.



*Figura 2.98 – Hub USB de sete portas.*

O barramento USB também evoluiu desde sua primeira versão (USB 1.1). O barramento USB original conseguia uma taxa de transferência de até **12 Mbps** (o equivalente a **1,5 MB/s**). O padrão USB 2.0 já é o mais comum nos atuais computadores (todos os computadores e equipamentos da atualidade são, em sua maioria, USB 2.0), e sua velocidade é de cerca de **480 Mbps** (isso mesmo! O equivalente a **60 MB/s** ou 40 vezes mais que o USB 1.1).

Já é possível, porém, encontrar barramento (e periféricos) USB 3.0. Essa nova geração de USB permite a transferência de informações a **4,8 Gbps** (10 vezes o USB 2.0!), o que equivale a **600 MB/s**.



*Figura 2.99 – Conector USB 3.0 (ligeiramente diferente do 2.0).*

Apesar de ligeiramente diferentes, os encaixes do USB 3.0 (portas na placa-mãe) aceitam a conexão de cabos USB das versões anteriores. O plug do tipo “A”, que é a parte do cabo que se encaixa justamente na placa-mãe, é muito semelhante aos anteriores. Já o plug “B”, que é a parte que encaixa no dispositivo periférico (impressora, scanner, disco rígido externo) é bem diferente, não sendo compatível com as versões antigas!



*Figura 2.100 – Cabo USB 3.0 – Plug “A” (em cima) e Plug “B”.*



*Figura 2.101 – Cabo USB 2.0 – Só para comparar!*

Assim como acontece no barramento serial ATA, o USB conta com a característica de ser **Hot Swap** (permitir a conexão e desconexão de dispositivos do computador sem precisar desligar ou reiniciar a máquina).

Além disso, o barramento USB é a verdadeira personificação do Plug and Play. Qualquer equipamento conectado a qualquer porta USB é automaticamente reconhecido pelo sistema operacional do computador, o que facilita muito a sua instalação.

Por ser dotado de Hot Swap e de Plug and Play, diz-se que o USB é um barramento **Hot Plug and Play**.

### **Barramento Firewire (IEEE 1394)**

Encontrado apenas em alguns computadores (em sua maioria laptops), o barramento firewire é incrivelmente rápido. Esse barramento foi criado originalmente para equipamentos de som, vídeo, instrumentos musicais e afins.

O barramento firewire foi regulamentado pela norma IEEE 1394. O IEEE é um instituto que reúne diversos cientistas e engenheiros em eletrônica e informática, que definem o funcionamento de diversos padrões da indústria mundial. Sua taxa de transferência atinge os 800 Mbps (ou 100 MB/s).





*Figura 2.102 – Porta firewire (IEEE 1394).*

Um único barramento firewire também pode ser usado por vários equipamentos ao mesmo tempo, em um total de 63 dispositivos. Também existem hubs firewire que funcionam de forma análoga aos hubs USB.

Por ser um pouco mais “caro” para a indústria, o IEEE 1394 não se tornou padrão de mercado, e provavelmente nem vai, especialmente com a chegada do USB 2.0 (e, agora, com o USB 3.0). Lembre-se: tanto o barramento USB quanto o firewire funcionam de forma SERIAL, ou seja, enviando um bit por vez.

### **Barramento Thunderbolt**

O barramento Thunderbolt foi criado pela Intel, em parceria com outras empresas, e está, aos poucos, sendo adotado pela indústria. Os micros Macintosh, da Apple, inclusive os laptops desta empresa, já trazem o barramento Thunderbolt consigo.



*Figura 2.103 – Porta Thunderbolt (a mesma da saída Minidisplay Port).*

Nos computadores da Apple®, a porta Thunderbolt é a mesma porta usada para conexão de monitores (Minidisplay Port), já vista no tópico sobre a placa de vídeo.

A conexão Thunderbolt promete entregar dados a uma velocidade de 10 Gbps (mais de 2x o que o USB 3.0 faz!) – é simplesmente surpreendente! Logo, logo, veremos mais conexões Thunderbolt nos computadores PC (já existem placas-mãe para PC com esse barramento, mas elas são, ainda, incomuns).

### **Bluetooth**

Usada para conexão de equipamentos sem uso de fios a curtas distâncias, a tecnologia bluetooth traz recursos muito interessantes. Com o bluetooth, praticamente qualquer equipamento seria ligado a um computador sem o uso de fios, através dos sinais de radiofrequência usados por essa tecnologia.

A tecnologia bluetooth permite que notebooks, micros de mesa, teclados, mouses, monitores, celulares, fones de ouvido e qualquer outro equipamento possam se comunicar apenas por ondas de rádio, ou seja, sem fios.

A ideia é que, quando um dispositivo equipado com bluetooth entra em uma área de cobertura da transmissão, ele é imediatamente localizado pelos demais equipamentos, e começa a se comunicar imediatamente. Ou seja, a tecnologia bluetooth é completamente plug and play (na verdade, hot plug and play, pois a detecção de um equipamento não requer o desligamento do computador).

A frequência de operação do bluetooth é de 2,4 GHz, e a distância ideal de conexão é de 10 metros. Um grande problema para essa tecnologia é que sua frequência pode sofrer interferências de outras frequências idênticas, como as praticadas pelos fornos de micro-ondas e alguns telefones sem fio, celulares e também as placas de rede sem fio da arquitetura 802.11.

Atualmente, embora já seja muito utilizado em informática, é comum usar o bluetooth para

realizar a comunicação entre dispositivos portáteis como celulares e tablets.

Bluetooth é uma tecnologia para criar WPAN (Wireless PAN – ou redes pessoais sem fio).

#### 2.4.8.3. RAID

Tecnologia para armazenamento de dados em HDs que aumenta os recursos do barramento utilizado (seja SATA ou SCSI). Com o **RAID** (Tabela Redundante de Discos Independentes) é possível **combinar vários HDs** para que estes funcionem como se fossem um único disco.

Quando ligamos dois HDs num computador, eles são completamente independentes entre si, “nem se ajudam, nem se atrapalham”, mas, com RAID, haverá uma relação estreita entre eles, que pode ser definida pelo modo de operação escolhido do RAID.

##### RAID 0 (Striping – Enfileiramento)

Dois ou mais discos rígidos funcionarão como um único, e suas capacidades serão somadas. Por exemplo, dois discos rígidos de 1 TB aparentarão ser um único disco de 2 TB. Os dados serão divididos entre os discos envolvidos, ou seja, quando um arquivo for gravado, metade dele vai para um disco, a outra metade é gravada no outro.

Isso torna o sistema muito mais rápido (tanto na leitura dos dados quanto na escrita destes), porque levará metade do tempo para se gravar um arquivo. Mas, nesse caso, se um dos discos falhar (pifar), não adiantará nada ter “meios-arquivos” no outro, e perderemos todos os dados.

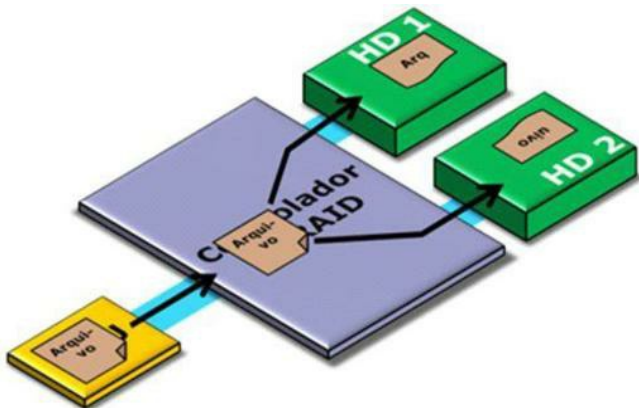


Figura 2.104 – RAID 0 com dois discos – cada disco grava parte do arquivo.

Resumindo... O RAID 0 traz velocidade, mas não segurança!

### RAID 1 (Mirroring – Espelhamento)

Dois ou mais discos rígidos funcionarão como apenas um, mas eles serão sempre cópias idênticas. Cada arquivo gravado é colocado em todos os discos ao mesmo tempo. Por exemplo, dois discos de 1 TB serão combinados e aparentarão ser, para o sistema operacional, um único disco de 1 TB.

O sistema fica mais rápido apenas no processo de leitura dos dados (em comparação a um sistema que não usa RAID). Não há ganho de velocidade no processo de escrita dos dados.

Embora não se garanta ganho de desempenho em todos os processos de uso do sistema de discos, certamente, haverá ganhos significativos no quesito de segurança do sistema.

No momento em que um disco rígido falhar, o outro assumirá imediatamente a sua posição sem que o sistema seja afetado, como se nada tivesse acontecido. Pouquíssimas vezes será necessário desligar e religar o computador. Isso confere ao RAID 1 a característica de **tolerância a falhas**, que o RAID 0 nem sonha em ter!

Pelo amor de Deus! Se você for fazer RAID no seu micro em casa, faça RAID 1. É bem mais seguro!

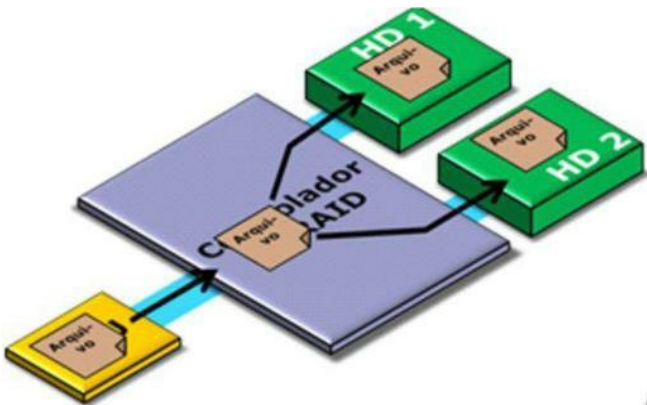


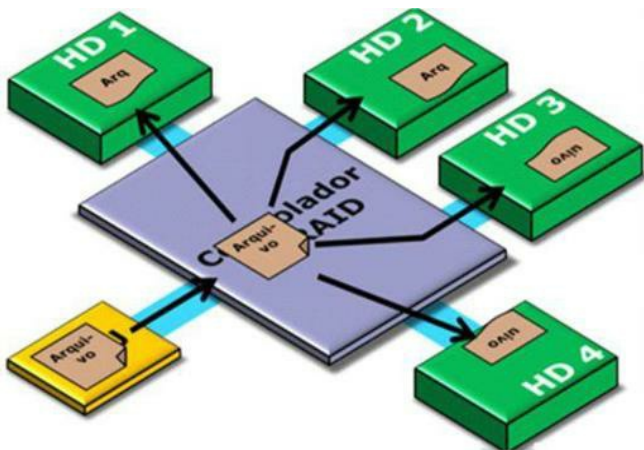
Figura 2.105 – RAID 1 com dois discos (o arquivo é gravado em ambos).

### **RAID 1+0 (também chamado RAID 10) (Mirror + Strip)**

Este modo de operação do RAID só pode ser executado com quatro discos rígidos (no mínimo).

Este é um modo combinação dos dois primeiros, em que uma dupla de HDs funcionará em RAID 0 (somando suas capacidades e acelerando o sistema) e os outros dois serão apenas uma cópia do primeiro par. Como você pode notar, essa opção é a mais cara de todas, pois exige a presença de quatro discos.

Na figura a seguir, os HDs 1 e 4 estão em RAID 0 (cada um deles está com uma “metade” do arquivo). Os HDs 2 e 3 também estão assim (em RAID 0). Agora é só entender que a dupla 1 (formada pelos HDs 1 e 4) está com o mesmo conteúdo (cópia) da dupla 2 (HDs 2 e 3) – logo, podemos concluir que as duas duplas estão em RAID 1.



*Figura 2.106 – RAID 10 – duas duplas de discos rígidos fazendo RAID 1 + 0.*

O RAID é muito comum em servidores de empresas para aumentar a velocidade e a confiabilidade do sistema, mas já está sendo possível aos usuários domésticos o acesso a esse tipo de tecnologia através das novas placas-mãe que incluem o RAID como recurso próprio. Praticamente todas as placas mãe de médio porte oferecem RAID nos barramentos SATA.

#### 2.4.9. Fonte de alimentação

Fonte de alimentação é um dispositivo que recebe a energia em corrente alternada da empresa elétrica (normalmente depois de passar por um estabilizador) e divide essa energia, distribuindo-a para os diversos dispositivos internos do gabinete, como a placa-mãe, a placa de vídeo, os HDs, entre outros.



*Figura 2.107 – Fonte de alimentação ATX.*

Uma das principais características das fontes cobradas em prova é o fato de as fontes atuais seguirem o padrão ATX de funcionamento (antigamente, lá pelos idos de 1998, as fontes eram consideradas apenas AT).

Uma fonte AT (antiga) era controlada manualmente, ou seja, via interruptores: funcionava quando o usuário apertava o interruptor para ligar o micro e parava de funcionar quando o usuário, manualmente, desligava aquele mesmo interruptor.

Lembra-se da época em que o Windows (95 e 98) esperava (com uma tela preta contendo a mensagem “SEU COMPUTADOR JÁ PODE SER DESLIGADO COM SEGURANÇA” em letras laranja) que o usuário desligasse o computador? Pois bem. Não era culpa do Windows, mas da fonte que não desligava automaticamente.

Desde o advento das fontes ATX, não temos tido o desprazer de ler tal mensagem. As fontes

ATX são eletrônicas. São controladas por software (leia-se, pelo sistema operacional, que sabe enviar mensagens à placa-mãe e esta as repassa à fonte). Com isso, as fontes ATX podem se desligar automaticamente, que é exatamente o que acontece hoje em dia, já que ao dar o comando ao Windows para desligar o micro, este repassa o comando à placa-mãe, que, por sua vez, repassa à fonte, que interrompe a energia para o gabinete, desligando-o automaticamente.

Uma fonte ATX também se liga automaticamente desde que o computador esteja configurado para “acordar” quando receber um estímulo específico (como uma chamada no modem ou uma comunicação via placa de rede). Para isso, o micro tem de estar em stand-by (modo de espera), ou seja, não estar totalmente desligado, apenas em estado de baixo consumo de energia.

Aliás, o recurso de estado de espera (baixo consumo) também só é possível graças às fontes ATX! Todo micro, hoje em dia, utiliza fontes no padrão ATX!

Outra coisa com relação às fontes é sua capacidade de carga (potência máxima que suporta). As fontes mais comuns suportam cerca de 300W de carga interna (1 placa-mãe, 1 HD, 1 gravador de DVD, 1 placa de vídeo normal). Aos que são mais “ambiciosos” no que se refere ao micro, talvez seja necessário adquirir alguma fonte mais potente (existem fontes de 500, 700 e até 1.000 Watts).

Certifique-se, porém, de ver se a potência da fonte é REAL. Apenas fontes de marcas muito boas e mais conceituadas (e mais caras!) são realmente potentes para gabinetes que contenham mais equipamentos.

## 2.5. Considerações finais sobre hardware

Bem, espero que os assuntos abordados neste capítulo tenham sido de grande valia para você, caro concursando, que busca incessantemente o conhecimento necessário para obter bons resultados em provas de informática.

Não sou, nem de longe, a pessoa mais conhecedora a respeito de hardware, mas resolvi compilar esse material com a máxima clareza para que não se torne ainda mais pesado o fardo que é estudar esse assunto.

Para mais informações a respeito de hardware, acesse:

<http://www.hardware.com.br>

Aproveito para parabenizar o autor do site, o professor Carlos E. Morimoto, por sua incrível facilidade em apresentar um conteúdo tão complexo para a maioria de nós, meros usuários. Com certeza, alguns dos melhores textos do professor Morimoto serviram de base para alguns assuntos aqui apresentados.

Para se manter atualizado com os novos lançamentos de equipamentos, novas tecnologias, testes e comparações entre marcas e modelos de dispositivos de informática, acesse:

<http://www.tomshardware.com> (em inglês)

ou

<http://www.clubedohardware.com.br>

ou ainda

<http://www.laercio.com.br>

Aqui vão minhas homenagens a dois monstros sagrados do estudo e ensino de hardware: Laércio Vasconcelos e Gabriel Torres. A vocês, professores, meu sincero agradecimento por

tornar, a cada dia, esse assunto tão complicado em algo fácil de digerir para meros mortais como nós.



## 2.6. Questões de hardware

Estilo FCC (Fundação Carlos Chagas)

1. (Técnico/TJ-PE/2007) O barramento especialmente desenvolvido para a comunicação da placa-mãe e a placa de vídeo é o:
  - a) PCMCIA;
  - b) PCI;
  - c) ISA;
  - d) AGP;
  - e) EISA.
2. (Técnico/TJ-PE/2007) A área de armazenamento temporário onde os dados frequentemente utilizados pelo processador são armazenados para acesso rápido é a:
  - a) ROM;
  - b) EDO;
  - c) CACHE;
  - d) SDRAM;
  - e) DDRAM.
3. (Técnico/TJ-PE/2007) Analise as seguintes afirmativas em relação aos chipsets das placas-mãe.
  - I. A ponte norte (north bridge) faz a comunicação do processador com as memórias e, em alguns casos, com os barramentos de alta velocidade.
  - II. Tidos como os principais circuitos integrados da placa-mãe, são responsáveis pelas comunicações entre o processador e os demais componentes.
  - III. A ponte sul (south bridge) é a responsável pelo controle de dispositivos de entrada ou saída (I/O), tais como interfaces IDE, drives de CD-ROM, de DVD-ROM e de disquete.É correto o que se afirma em:
  - a) I, II e III;
  - b) II e III, apenas;
  - c) I e II, apenas;
  - d) III, apenas;
  - e) I, apenas.
4. (Técnico/TJ-PE/2007) Em relação a slots de conexão, é correto afirmar:
  - a) placas PCI mais velozes devem ser instaladas nos slots PCI mais próximos do processador;
  - b) quando uma placa-mãe não tem slot AGP, a única opção de uso é o vídeo onboard;
  - c) o slot AGP 1x tem taxa de transferência menor que o slot PCI, que, por sua vez, tem taxa de transferência menor que o AGP 2x;
  - d) um slot AGP 8x é mais veloz que um slot PCI Express 16x;
  - e) slots PCI Express têm velocidade superior aos slots AGP.

5. (Técnico/TJ-PE/2007) Um disco de capacidade nominal de 40 GB (informado pelo fabricante) está formatado com partição única, no Windows XP. Ao se clicar nas propriedades desse disco o Windows XP indica que o disco tem uma capacidade menor. Isso acontece porque o:
- a) fabricante utilizou como medição múltiplos de 1.000 e não 1.024 bytes, que daria o total de 42.949.672 bytes, equivalentes a 40 GB;
  - b) disco formatado com FAT32 tem uma parte do seu espaço desperdiçado;
  - c) disco consumiu a diferença na instalação do sistema operacional;
  - d) BIOS da placa-mãe só reconheceu 74,5 GB;
  - e) disco não foi corretamente formatado.

Estilo ESAF

1. (AFC/CGU/2004) Analise as seguintes afirmações relativas a componentes de hardware de computadores.
- I. A placa-mãe é a principal placa de circuitos de um microcomputador. O único componente que não pode ser instalado ou equipar uma placa-mãe é o barramento AGP.
  - II. O barramento AGP é o primeiro barramento a possuir um slot que permite expansão, opera com 8 bits e em sua segunda versão, ampliada e melhorada, opera com 16 bits para dados e 24 bits para endereçamento, com uma frequência de operação de 8MHz.
  - III. Uma característica importante dos dispositivos PCI é o Plug and Play. Esses dispositivos são equipados com uma memória ROM contendo informações que permitem ao sistema operacional detectá-los automaticamente.
  - IV. Um computador, alimentado por uma fonte com padrão ATX e com uma placa-mãe apropriada para esse padrão, permite que seja ligado ao receber um sinal externo como, por exemplo, uma chamada telefônica recebida pelo modem nele instalado.
- Estão corretos os itens:
- a) I e II;
  - b) II e III;
  - c) III e IV;
  - d) I e III;
  - e) II e IV.
2. (AFRF/TI/2005) Analise as seguintes afirmações, relacionadas aos componentes funcionais (hardware) de um computador.
- I. Em uma placa-mãe, as entradas padrão PCI servem para encaixar os cabos que ligam unidades de CD/DVD. Esses cabos, chamados de flat cables, podem ser de 40 ou 80 vias. Cada cabo pode suportar até duas unidades de CD/DVD.
  - II. O endereçamento consiste na capacidade do processador de acessar um número máximo de células da memória. Para acessar uma célula, o processador precisa saber o endereço dela. Cada célula armazena um byte. Assim, um processador com o barramento de dados com 16 bits pode acessar duas células por vez.
  - III. O clock interno indica a frequência na qual o processador trabalha. Portanto, num

Pentium IV de 2,6 GHz, o “2,6 GHz” indica o clock interno, geralmente obtido por meio de um multiplicador do clock externo. O clock externo é o que indica a frequência de trabalho do barramento de comunicação com a placa-mãe.

IV. O setor de BOOT de um HD contém um pequeno software chamado Post, que é responsável por controlar o uso do hardware do computador, manter as informações relativas à hora e à data e testar os componentes de hardware após o computador ser ligado. Indique a opção que contenha todas as afirmações verdadeiras.

- a) I e II.
- b) II e IV.
- c) III e IV.
- d) I e III.
- e) II e III.

3. (AFRE/TI/2005) Com relação à arquitetura de computadores é correto afirmar que:

- a) a arquitetura RISC especifica que o microprocessador possui poucas instruções, mas cada uma delas é otimizada para que sejam executadas muito rapidamente, normalmente, dentro de um único ciclo de relógio;
- b) o BIOS é o circuito de apoio ao computador que gerencia praticamente todo o funcionamento da placa-mãe (controle de memória cache, DRAM, controle do buffer de dados, interface com a CPU etc.). Ele é responsável pelas informações necessárias ao reconhecimento de hardware (armazenadas na sua memória ROM);
- c) usando-se um endereço de K bits, pode-se endereçar no máximo  $K^2$  ( $K \times K$ ) posições de memória ou células de memória;
- d) o chipset é um pequeno programa armazenado na memória ROM da placa-mãe. É responsável por acordar o computador, contar e verificar a memória RAM, inicializar dispositivos, e o principal, dar início ao processo de boot;
- e) os registradores são memórias ROM utilizadas para o armazenamento de dados.

4. (AFRE/2002) Em um computador, o objetivo do barramento é reduzir o número de interconexões entre a CPU e seus subsistemas. Para evitar a necessidade de um elevado número de caminhos de comunicação entre a memória e cada um dos dispositivos de entrada e saída, a CPU é interconectada com sua memória e sistemas de entrada e saída via um barramento de sistema compartilhado. Com relação à funcionalidade dos barramentos e acessos à memória em um computador é correto afirmar que:

- a) a memória gera endereços que são colocados no barramento de endereços, e a CPU recebe endereços do barramento de endereços;
- b) a CPU e a memória geram endereços que são colocados no barramento de endereços, e a memória recebe endereços do barramento de endereços;
- c) a CPU gera endereços que são colocados no barramento de endereços, e a memória recebe endereços do barramento de endereços;
- d) a CPU gera endereços que são colocados no barramento de endereços, e a CPU e a memória receberão endereços do barramento de endereços.
- e) tanto a CPU quanto a memória geram endereços que são colocados no barramento de

endereços e recebem endereços do barramento de endereços.

#### 3.1. Pequena definição sobre software

É correto afirmar que o computador é um conjunto de componentes eletrônicos que trabalham de forma harmoniosa para processamento de informações. Mas não é só isso... Os componentes físicos do computador (hardware) são simples peças do jogo do processamento. Quem as comanda são os **softwares** (programas). Programas de computador, ou softwares, são instruções digitais, gravadas em um computador, executadas pela CPU do computador no momento devido.

Um exemplo simples: um jovem acorda, chega à frente da porta da geladeira e se depara com um bilhete escrito por sua irmã mais velha:

1. Limpe seu quarto;
2. Lave o carro;
3. Compre ovos e leite;
4. Leve o cachorro para passear.

O que é isso?

“É uma exploração, João!”

Sem mencionar a exploração, caro leitor, é um conjunto de ordens a que ele deve obedecer. Isso é um exemplo de programa. Só que as “ordens” são dadas ao computador, que as executa sem questionar nem se revoltar.

Se o rapaz decidir aceitar a programação que lhe é imposta, aqui vão as comparações:

- **Bilhete:** programa (software);
- **Rapaz:** CPU (quem irá executar as tarefas);
- **Irmã:** programador (quem escreve o roteiro a ser seguido pela CPU);
- **Porta da geladeira:** memória auxiliar (onde o programa fica gravado até ser executado);
- **Memória do rapaz:** memória principal (onde o programa se manterá enquanto estiver em execução).

#### 3.2. Como funciona um programa?

Todo programa de computador é criado por alguém (o programador), e, de alguma forma, chega até o computador da pessoa que irá utilizá-lo (usuário). Enquanto estiver “dormindo”, sem ser usado, um programa está gravado em uma memória auxiliar (normalmente o disco rígido) na forma de **arquivos**.

Quando o programa entra em execução, ou seja, quando começa a funcionar, seus dados (ou parte deles) são copiados para a memória RAM, de onde são requisitados pela CPU durante todo o processo de execução.

Nota: os dados são copiados, pois mesmo quando estão em execução, com suas informações na RAM do computador, os programas continuam existindo nos arquivos gravados no disco rígido. Se alguma questão falar em “os dados de programas são **movidos** para RAM durante a execução destes programas”, a resposta é **“ERRADO”**.

Veja um programa muito comum no Windows: a **calculadora**. Ela é apenas um conjunto de instruções binárias gravadas no disco rígido do computador. Quando o usuário solicita o início de sua execução, através do clique na opção Calculadora dentro do menu Acessórios, o programa é imediatamente copiado para a memória RAM, de onde seus dados e instruções são buscados pela CPU durante o funcionamento do programa.



*Figura 3.1 – Calculadora em execução: dados e instruções na memória RAM.*

Sim, mas qual a diferença entre dados e instruções? Instruções são as ordens que estão no programa e que a CPU tem de executar. Como as quatro ordens no bilhete anterior.

“E as instruções dos programas têm de ser escritas naquele conjunto de instruções (“idioma”) que a CPU entende, não é, João?”

Sim, precisamente, leitor! O conjunto de instruções que a CPU entende é justamente importante porque qualquer programa que se diz feito para aquela CPU é construído especificamente com aquele conjunto de instruções.

Dados são as informações obtidas pelo programa, ou fornecidas a ele, para que haja funcionamento correto (por exemplo, o leite e os ovos, que serão trazidos; ou o cachorro, que será levado para o passeio).

### 3.3. Tipos de softwares

Há vários tipos de softwares disponíveis no mercado, cada um com uma finalidade, mas que podem ser divididos nestas categorias:

1. Software básico
  - 1.1. Sistemas operacionais
  - 1.2. Linguagens de programação
  - 1.3. Tradutores (compiladores/interpretadores)
2. Softwares utilitários
3. Softwares aplicativos

Conheça-os mais:

Sistemas operacionais são softwares que gerenciam os recursos do computador, fazendo-o funcionar corretamente; sem um sistema operacional, o computador não funcionaria. Vamos abordar o tema mais à frente.

Linguagens de programação são os códigos usados pelos programadores para criar os softwares. São os “idiomas” de alto nível que os programadores (criadores de programas) usam para escrever os códigos que darão origem aos seus programas.

Tradutores são os softwares responsáveis por transformar o código criado pela linguagem de programação (alto nível) em software executável (programa na linguagem das instruções da máquina) propriamente dito.

“Por que há a necessidade de tradução, João?”

Simples, caro leitor! Porque as linguagens de programação (chamadas de “alto nível”) são entendidas pelos programadores (são mais “próximas” da nossa linguagem – inglês), mas não pelas máquinas. As máquinas só entendem 0 (zero) e 1 (um) devidamente organizados para serem interpretados como instruções.

Ou seja, a maioria dos programadores não fala (nem escreve) diretamente na linguagem da máquina. Eles escrevem em linguagens mais “acessíveis” e depois traduzem os códigos que criam na linguagem da máquina!

Utilitários são programas que permitem a manutenção dos recursos da máquina, como ajustes em discos, memória, conserto de outros programas etc. Ex.: antivírus, programas de melhoria de desempenho, gerenciamento e aproveitamento de memória, entre outros.

Aplicativos são softwares voltados para a solução de problemas dos usuários, como os programas para planilhas de cálculos, edição de texto, desenho, bancos de dados, edição de fotos etc.

Seja qual for o tipo do software, ele é um conjunto de instruções binárias gravadas em uma memória permanente na forma de um ou mais arquivos.

### 3.4. O que são arquivos?

Todas as informações, quer sejam instruções de programas, quer sejam dados, são gravadas em memórias, como vimos nas aulas de hardware (Capítulo 2).

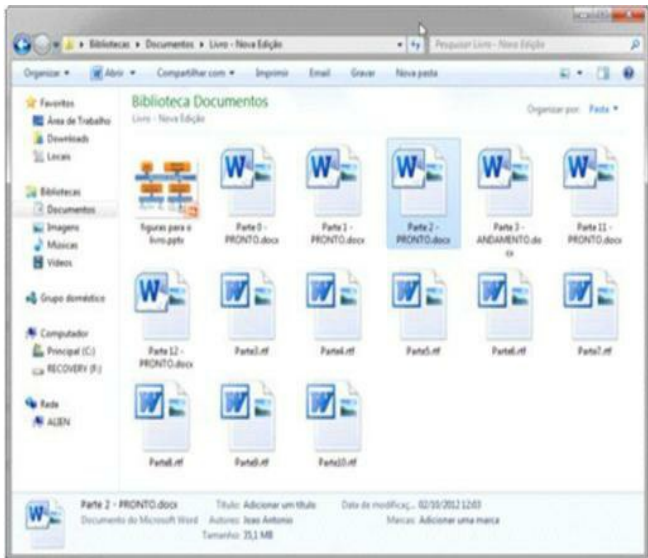
Essas informações podem ser gravadas em memórias permanentes, chamadas unidades de armazenamento, que são normalmente memórias de disco (como os pen drives, HDs e DVDs). Quando gravadas em qualquer uma dessas memórias permanentes, as informações são reunidas

em blocos ordenados definidos, chamados **arquivos**.

Arquivo pode ser definido de várias formas, e com números variados de verbetes, mas gosto de defini-lo assim: **arquivo é um bloco de informações relacionadas, que está gravado em uma unidade de armazenamento**. Um arquivo tem de ter um nome, para que se possa identificá-lo e diferenciá-lo dos demais arquivos na mesma unidade.

Tudo o que “salvamos” no computador vira arquivo. Todos os programas em nossa máquina são gravados na forma de arquivos. Todas as informações que temos o direito de acessar no nosso computador ou em qualquer computador da Internet são arquivos.

Um simples exemplo: ao digitar um endereço de Internet qualquer, como <http://www.qualquercoisa.com/apostilas/testes.pdf>, na verdade você está “solicitando” o arquivo testes.pdf que está localizado no computador denominado www.qualquercoisa.com.



*Figura 3.2 – Vários arquivos.*



### **3.5. O que são pastas?**

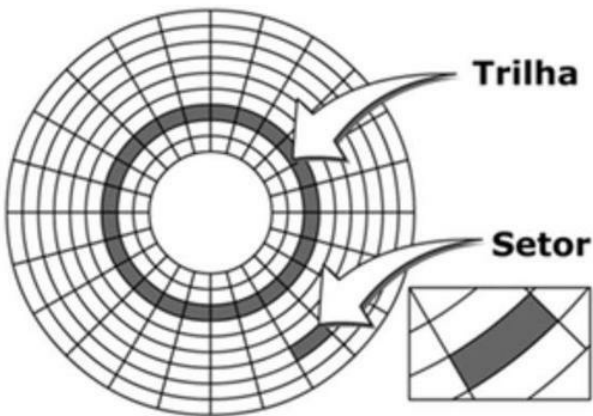
São pequenos compartimentos lógicos, criados em uma unidade para organizar melhor seu conteúdo para o usuário. Pastas, também conhecidas como diretórios, são meramente “gavetas” que podem guardar arquivos ou outras pastas.

As pastas não são informação importante para o usuário, ou seja, não há imagens, textos, sons numa pasta. Elas são simplesmente “cômodos” para armazenar os arquivos visando a mais rápida localização, por parte do usuário. Imagens, sons, textos, planilhas são, na verdade, arquivos que, por sua vez, ficam armazenados nas pastas que criamos para nos organizar.

### **3.6. Estrutura dos discos**

De novo: as informações digitais, quer sejam programas, quer sejam dados do usuário, são gravadas em unidades de armazenamento devido ao fato de essas unidades serem memórias permanentes. É interessante conhecer onde e como, exatamente, essas informações binárias são gravadas na superfície das unidades de armazenamento, como os HDs e disquetes.

Tomemos como exemplo um simples disquete (embora o exemplo sirva perfeitamente para HDs e CDs também): o disquete é uma memória em forma de disco, com superfície de gravação magnética (neste caso, há diferenças em relação ao CD, que não é magnético e usa superfície óptica de gravação), que é dividida em círculos concêntricos chamados trilhas. Essas trilhas, por sua vez, são divididas em pequenas unidades para armazenamento, chamadas setores. Veja a figura a seguir.



*Figura 3.3 – A estrutura física da superfície de uma unidade de disco (disquete, HD, CD).*

Os setores são, efetivamente, os locais onde os dados digitais são armazenados. Um setor possui, em uma unidade magnética como o HD ou o disquete, uma capacidade de armazenamento de 512 bytes. O setor de um CD possui uma capacidade diferente do setor das unidades magnéticas: 2.048 bytes.

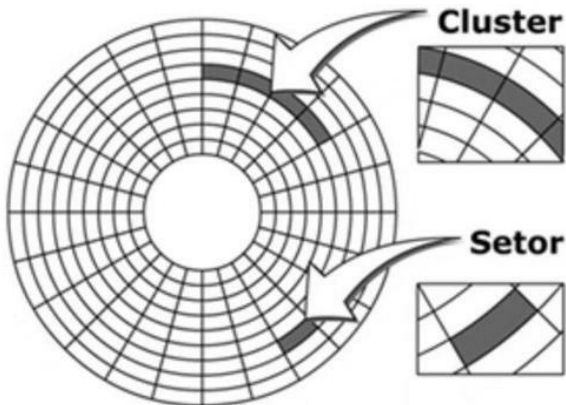
Mas, infelizmente, não são todos os sistemas operacionais que conseguem entender a grande quantidade de setores que há num disco rígido (depende do sistema operacional); portanto, os setores são reunidos em pequenos grupos chamados **clusters**, que passam a ser a mínima quantidade de informação que um disco consegue entender.

Um cluster é um conjunto de setores contíguos, que são reunidos simplesmente para que seja possível gerenciar o conteúdo do disco. Como um disco rígido tem muitos setores, gerenciá-los um a um seria muito custoso para o sistema operacional. Pense em uma escola com mais de mil alunos tentando arrumar o horário de todos eles individualmente (aulas particulares para todos). Seria inviável propor aulas particulares para todos porque seria impossível gerenciar todos os alunos como entidades individuais.

Por isso os alunos são reunidos em classes (turmas), o que, para a direção da escola, facilita muito a vida na hora de definir horários de aulas. Afinal, é mais fácil gerenciar os horários de 20 turmas do que de 1.000 alunos, não acha?

Nos discos rígidos, um cluster é, portanto, uma reunião necessária de setores visando ao

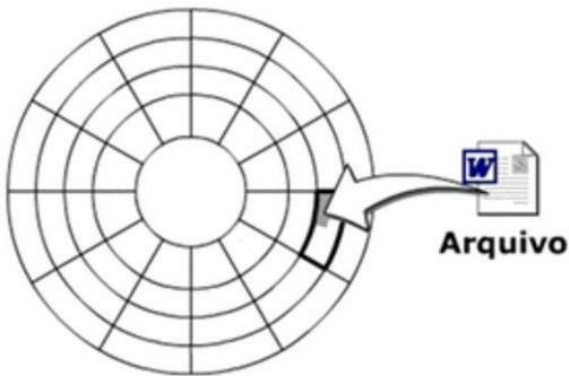
perfeito gerenciamento dos dados armazenados no disco.



*Figura 3.4 – Note que um cluster é, na verdade, apenas um conjunto de setores.*

Lembre-se: cluster é a menor quantidade de informação que um sistema operacional consegue gerenciar em um disco. Normalmente, um cluster é formado por vários setores, e mesmo que um arquivo ocupe apenas metade de um setor, ele será considerado como se ocupasse, no disco, um cluster inteiro. E aquele cluster será utilizado de forma única, por apenas um arquivo, quer dizer que não pode haver dois arquivos guardados no mesmo cluster.

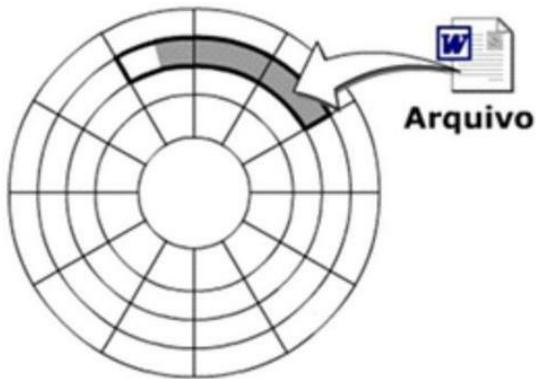
Ou seja, um cluster é a **menor unidade de alocação** de arquivos em um disco.



*Figura 3.5 – Um arquivo ocupando parte de um cluster: o restante do cluster é considerado ocupado.*

Diante disso, quanto maior o tamanho do cluster, maior é o desperdício de espaço no disco em questão. Funciona como a ideia de consumação mínima: um arquivo pode ter apenas 5 KB, mas se o tamanho do cluster é de 32 KB, o arquivo ocupará esse cluster para si, desperdiçando 27 KB. Se o arquivo ocupa 19 KB, vai desperdiçar 13 KB e assim por diante. Mesmo que o arquivo seja muito pequeno (50 bytes), ele será armazenado em um cluster inteiro, o que o faz ocupar 32 KB de espaço no disco (isso, é claro, levando em conta que o cluster tem esse tamanho).

Mas, atenção: um arquivo pode ocupar mais de um cluster, dependendo da quantidade de informações que ele possui e da capacidade de armazenamento do cluster.

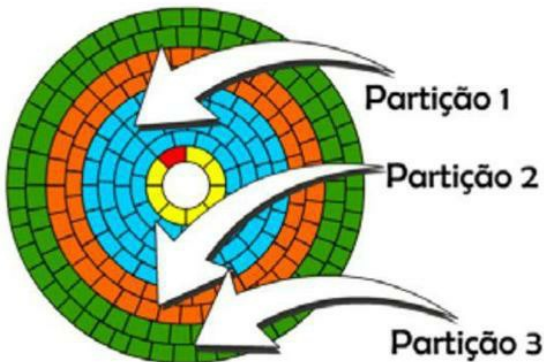


*Figura 3.6 – Um arquivo ocupando vários clusters.*

Lembre-se: trilhas e setores são características físicas do HD, mas os clusters são lógicos, definidos pelo sistema de arquivos utilizado.

Há ainda outra divisão de um disco (essa acontece somente em discos rígidos, não sendo usada em nenhum outro tipo de disco): as **partições**. Um mesmo disco rígido pode ser dividido em algumas seções (as partições), que funcionarão como se fossem discos diferentes para o sistema operacional.

Uma partição é uma divisão lógica do disco rígido. Cada partição será vista como uma unidade independente pelo sistema operacional.



*Figura 3.7 – Disco dividido em três partições.*

Entendendo a figura anterior: as partições não são divisões que fisicamente apresentam-se como “fatias”, mas na forma de coroa circular (conjuntos de trilhas). Então, as partições seriam anéis (um após o outro) a partir do centro do disco! (É só para explicar melhor.)

Um único disco rígido pode possuir diversas partições de tamanhos variados. As partições terão seu próprio número estipulado de clusters dependendo da divisão ocorrida no disco. O particionamento é muito comum quando se deseja instalar mais de um sistema operacional no mesmo computador.

As características acerca dos clusters (como quantidade de setores que os formam, quantidade máxima de clusters por partição, forma de identificação e indexação dos mesmos) variam muito de acordo com o sistema operacional usado (DOS, Windows XP, Linux, Windows 7 etc.), porém, mais precisamente, de acordo com o *sistema de arquivos* utilizado.

### 3.7. Sistema de arquivos

Um sistema de arquivos é um conjunto de rotinas (regras) que um determinado sistema operacional deve seguir para acessar unidades de disco, tanto na hora de gravar informações, quanto quando as lê.

Quando definimos que um disco será escrito neste ou naquele sistema de arquivos? Existem programas utilitários que criam as estruturas lógicas dos discos e os formatam (formatar seria “preparar para o uso”) com um devido sistema de arquivos desejado pelo usuário e entendido pelo sistema operacional em questão. Portanto, o sistema de arquivos que uma partição vai usar é

definido durante a formatação daquela partição.

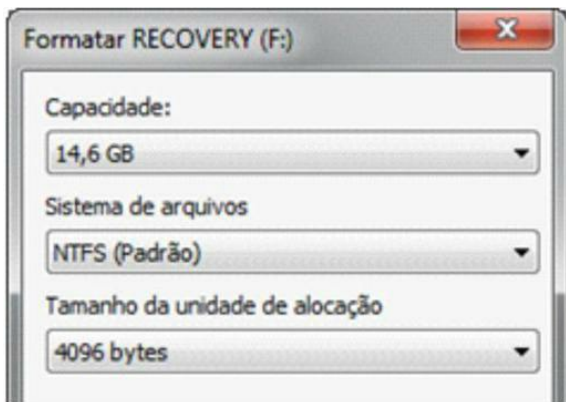


Figura 3.8 – Tela de formatação – escolha do sistema de arquivos da partição.

Diferentes sistemas operacionais fazem uso de diferentes sistemas de arquivos, pois cada sistema de arquivos é fabricado quase que exclusivamente para um determinado sistema operacional.

**Atenção:** um disco rígido com várias partições pode ter um sistema de arquivos em cada uma delas. Isso é muito comum quando se instalam dois sistemas operacionais que usam sistemas de arquivos diferentes no mesmo computador.

Os principais sistemas de arquivos vistos em concursos são os sistemas usados pela família Microsoft® de sistemas operacionais (DOS, Windows). São eles: FAT16, FAT32 e NTFS. O sistema operacional Linux usa seus próprios sistemas de arquivos (EXT2, EXT3, Reiser, entre outros).

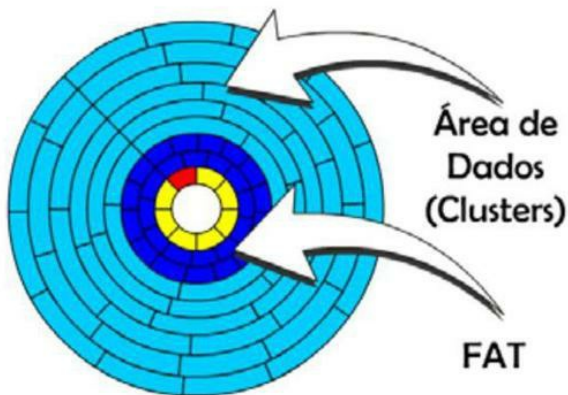
### 3.7.1. FAT (Tabela de Alocação de Arquivos)

A FAT nada mais é que uma espécie de índice (ou mapa, se preferir) gravado no início do disco (nas primeiras trilhas) para localizar com precisão todos os clusters existentes no disco.

Quando gravamos um arquivo, seja do Word, do Excel, ou de qualquer programa, ele é armazenado pelo Windows (o sistema operacional) em um (ou mais de um, dependendo do tamanho do arquivo) cluster disponível do disco, e sua posição inicial é gravada na FAT para que

possa ser encontrado posteriormente.

Todas as vezes que solicitamos a abertura de um arquivo (por exemplo, quando damos duplo clique em algum ícone), o sistema operacional localiza, através da FAT, a posição correta do arquivo no disco e começa a lê-lo. Veja a figura a seguir.



*Figura 3.9 – A FAT (Tabela de Alocação) é o índice para achar os arquivos gravados no disco.*

Quando apagamos um arquivo do disco (mesmo depois de limpar a lixeira), apenas estamos tirando a informação da FAT, com isso ele não poderá mais ser encontrado pelo sistema operacional.

Quando formatamos um disco (supostamente significa apagar todos os dados dele), na verdade, estamos limpando o conteúdo da FAT, o que, aparentemente, mostrará o disco vazio. (Porém, as informações binárias dos arquivos estão lá, só não estarão acessíveis para o sistema operacional!)

Atenção: existem programas que conseguem recuperar um arquivo através da leitura de seus dados em seus clusters, ignorando que a FAT informe que o arquivo não existe. São utilitários muito interessantes. É com o uso desses programas, por exemplo, que a Polícia Federal, a Secretaria de Fazenda, entre outros órgãos e instituições, conseguem recuperar arquivos supostamente deletados (excluídos) de computadores suspeitos.

A “família” dos sistemas de arquivos que usam a FAT é composta por:



- **FAT16:** sistema de arquivos antigo, usado no DOS e Windows 95, mas também suportado pelos Windows mais novos. Usava clusters de até 32 KB e só conseguia gerenciar até cerca de 65 mil clusters (65.536 clusters para ser exato). O que totaliza 2 GB de tamanho máximo para a partição com FAT16.

Caso você use um pen drive de capacidade maior (16 GB, por exemplo), a FAT16 só entenderá 2 GB do disco, os outros 14 GB não serão reconhecidos, serão simplesmente perdidos.

Como todos os discos rígidos atuais e a maioria dos pen drives da atualidade apresentam capacidades maiores que essa, o FAT16 é simplesmente inviável. O FAT16 também apresentava uma limitação: o nome dos arquivos tinha de ter, no máximo, 11 letras (oito para o nome e três para a extensão). Um arquivo não poderia se chamar *Relatório de Vendas.doc*, mas sim algo como *relvenda.doc*.

- **FAT32:** Uma evolução natural do FAT16, o FAT32 tem mais recursos que o antecessor, como o aumento da capacidade máxima gerenciável da partição. Uma partição formatada com FAT32 pode ter, no máximo, 2.048 Gigabytes de capacidade (o que equivale a 2 TB – já chegamos a isso em alguns HDs).

Outra limitação interessante do FAT32 é o tamanho máximo do arquivo que o sistema consegue reconhecer: 4 GB. Ou seja, ao salvar um arquivo em FAT32, esse arquivo não pode ter mais de 4 GB (muitos vídeos atuais e arquivos compactados de backup já atingem esse valor).

Logo, FAT32 não é recomendado para quem trabalha editando vídeos, por exemplo, já que podem, facilmente, atingir o tamanho de 4 GB para os arquivos que manipulam.

### 3.7.2. NTFS – Sistema de arquivos do NT

Sistema de arquivos desenvolvido pela Microsoft® para os sistemas operacionais corporativos (Windows NT, Windows 2000, Windows 2008), além dos mais recentes sistemas Windows domésticos e de uso geral: Windows XP, Windows Vista e o Windows 7. O Windows 8, com lançamento previsto para 2013, também entenderá NTFS, claro!

NTFS significa NT File System, e é uma excelente forma de gravar arquivos em um disco, apresentando uma série de vantagens em relação às partições FAT16 e FAT32. Nesse sistema de arquivos os clusters máximos são de apenas 4 KB, o que garante um melhor aproveitamento das partições, qualquer que seja o tamanho delas.

Não existe a FAT no sistema NTFS, mas o “índice” dos clusters se chama MFT (Tabela Mestre de Arquivos), que funciona muito melhor que a FAT, apresentando diferenças significativas.

Enquanto a FAT apenas aponta para o cluster inicial do arquivo, informando ao sistema operacional em que posição do disco ele está, a MFT guarda dentro de si informações básicas sobre o arquivo, além de um pequeno trecho dos seus dados. (Isso mesmo, o início do arquivo é gravado na própria MFT.)

Segue uma lista das principais características do NTFS:

- **Segurança de Acesso:** através dos recursos disponibilizados pela MFT é possível definir níveis de acesso aos arquivos gravados na partição formatada com NTFS. Desse modo, um arquivo só será acessado por quem realmente tiver autorização para tanto.
- **Cota de Disco:** é possível definir limites de tamanho de armazenamento para os usuários do sistema (ou seja, tal usuário só terá direito de armazenar até 1 GB de dados naquele dado

disco).

- **Criptografia:** embaralhamento automático dos arquivos e pastas gravados para que não possam ser identificados e/ou lidos por pessoas não autorizadas (mesmo que se roube aquele HD).

- **Compactação:** arquivos e pastas podem ser automaticamente compactados (reescritos de forma que ocupem menos espaço no disco).

- **Clusters Personalizados:** o usuário tem o direito de escolher o tamanho dos clusters de uma partição formatada com NTFS, não importando qual seja o tamanho da partição em si.

Outra característica que torna o NTFS muito melhor que os FAT (especialmente o FAT32, que ainda é o mais usado hoje) é que não há a limitação de um arquivo com 4 GB nem de uma partição de 2 TB.

No NTFS, podemos ter arquivos de 16 EB (Exabytes) e partições também de 16 EB. (Totalmente absurdo, não?!)

Tanto o NTFS quanto os sistemas FAT podem ser usados, além dos HDs, em pen drives, SSDs e cartões de memória. CDs e DVDs usam seus próprios sistemas de arquivos, que não são tão interessantes de estudarmos.

### 3.8. Processo de inicialização do computador

Quando ligamos o nosso computador, presenciamos uma série de acontecimentos ordenados e previamente programados para que o computador possa funcionar corretamente e nos permita utilizá-lo.

O conjunto desses acontecimentos é denominado **boot** (ou processo de inicialização). O boot inicia-se com a ligação propriamente dita do computador, que passa a receber energia elétrica para sua alimentação. Após o computador ser ligado, um pequeno programa chamado **BIOS** (Sistema Básico de Entrada e Saída) é executado, dando início a algumas operações que estão determinadas em seu roteiro.

O BIOS está permanentemente gravado em um chip de memória ROM que fica localizado na placa-mãe do computador. Atualmente, o BIOS é gravado em um CHIP de memória EEPROM ou memória flash, que permite sua alteração (normalmente necessária para uma atualização desse programa).



*Figura 3.10* – BIOS, o responsável por “dar a partida” no microcomputador.

Lembre-se: se perguntarem se o BIOS pode ser “atualizado”, a resposta é SIM! Mas isso só acontece se este programa for gravado em um chip de memória EEPROM ou FLASH (variantes “alteráveis” da ROM). Se a pergunta for, porém: “O BIOS, em memória ROM, pode ser atualizado?” – neste caso, NÃO, porque a ROM é imutável (lembre-se disso!).

Ao ser iniciado, o programa BIOS é carregado para a RAM, onde é efetivamente executado, e realiza uma checagem de rotina para verificar quais são os equipamentos e componentes ligados ao computador. O BIOS verifica se há processador, conta a memória RAM, localiza o HD, teclado, monitor, placa de vídeo, placas nos slots ISA, PCI, AGP, entre outros.

A checagem que o BIOS realiza é chamada **POST** (Power On Self Test – algo como “autoexame na ligação”) e visa fornecer uma descrição completa do seu computador para o sistema operacional.

Não é incomum, porém, questões de prova em que POST é “o programa responsável por detectar o hardware do computador” – mesmo sendo apenas “parte de um programa” (uma parte das instruções presentes no programa maior: o BIOS).

Depois de concluído o POST, o BIOS busca imediatamente o sistema operacional (qualquer que seja ele, Windows, Linux etc.) e então carrega suas informações para a RAM. A partir desse momento, o BIOS volta a “dormir” e entrega o trabalho para o sistema operacional. A sequência é mais ou menos esta:

1. Usuário liga o computador;
2. BIOS é carregado para a RAM e realiza as primeiras operações;
3. BIOS realiza o POST (checagem dos componentes básicos);
4. BIOS procura o sistema operacional numa memória auxiliar;
5. Sistema operacional assume o controle da máquina.

Certo, tudo bem... Entendido até agora? Mas, exatamente onde o sistema operacional deve estar para que o BIOS o encontre?

Como é possível que todas as vezes em que ligamos o computador, o sistema operacional (no nosso caso, o Windows) seja iniciado normalmente? Como o BIOS sabia onde ele estava? Resposta: o BIOS foi procurar o sistema operacional no HD (disco rígido), mais precisamente no setor de boot do HD.

### 3.8.1. Setor de boot

Todos os discos (HDs, CDs, disquetes, entre outros) possuem um pequeno setor reservado para as informações básicas de programas que irão carregar o sistema operacional. Esse setor é conhecido como **setor de Boot**.

Consiste em apenas um único setor do disco (512 bytes, por exemplo, no HD) que armazena um pequeno programa chamado carregador do sistema operacional (boot loader). As instruções desse programa foram desenvolvidas para encontrar e carregar (na memória) os verdadeiros arquivos que formam o sistema operacional naquele computador.

Em um HD, o setor de boot é chamado, muitas vezes, de **MBR (Master Boot Record – Registro Mestre de Inicialização)**, mas, atenção, esse nome só vale para o HD!

“Mas, João, setor de boot e MBR são a mesma coisa?”

Não, caro leitor, o MBR é um espaço inicial num HD e guarda, além do setor de boot, outro componente que só existe nos HDs: a **tabela de partições**.

Uma tabela de partição é um espaço de memória dentro do MBR que registra quantas são as partições existentes nos discos. Uma tabela de partições só consegue registrar a existência de quatro partições.

Daí a razão de somente um disco rígido poder ser dividido em partições: só há tabela de partições nos discos rígidos. Todos os discos têm setor de boot, mas só os HDs contêm MBR, e, logo, só os HDs possuem tabelas de partições.

Portanto, caro leitor, é possível dizer que: **MBR = setor de boot + tabela de partições**.

Independentemente do sistema operacional utilizado, um apontador para suas informações iniciais é gravado no MBR para que o BIOS do computador sempre consiga encontrá-lo. Como o Windows está instalado em nosso disco rígido, o BIOS localiza imediatamente o MBR do HD e encontra o que necessita (o programa carregador do Windows) para chamar o Windows para o trabalho.

Pode haver casos em que o sistema operacional não seja iniciado pelo HD, por razões diversas

(incluindo alguns vírus que têm o péssimo hábito de apagar o conteúdo do MBR); assim sendo, faz-se necessária a existência de um disquete de boot, ou disquete de inicialização. Hoje em dia, claro, é mais comum que haja CDs ou DVDs de boot.

Também é possível, devido aos seus tamanhos atuais, possuir um “pen drive de boot”, ou um “pen drive inicializável”, que contém o sistema operacional em sua memória e que possui seu setor de boot devidamente registrado, contendo o carregador daquele sistema.

Um pen drive como esses, quando for colocado em uma porta USB de um computador atual desligado, será usado para carregar seu sistema diretamente para a RAM daquele computador, permitindo que o computador seja ligado com aquele sistema operacional sem utilizar o sistema já instalado no HD daquela máquina. (Muita gente usa o Linux assim!)

### **3.8.2. Múltiplos sistemas operacionais**

É possível instalar em um mesmo computador vários sistemas operacionais, mas o usuário deverá, durante o boot, escolher qual usará naquele momento. Para se ter vários sistemas operacionais, deve-se particionar o HD em várias unidades (normalmente uma para cada sistema operacional a ser instalado).

Como o MBR é muito pequeno, não dá pra armazenar as informações dos dois sistemas operacionais ao mesmo tempo; portanto, no caso de um DUAL BOOT (dois sistemas operacionais), por exemplo, é gravado no MBR um programa chamado boot manager (gerenciador de boot). Esse programa irá solicitar que o usuário escolha qual sistema operacional será utilizado, e, quando for escolhido, o programa tratará de inicializar o sistema operacional desejado.

Note, na figura seguinte, que o programa oferece algumas opções de sistemas operacionais no computador, como o Linux e o Windows. A foto seguinte é do programa LILO (distribuído com o Linux). Há outro gerenciador de boot que acompanha o Linux: o GRUB.

GNU GRUB version 0.95 (638K lower / 186384K upper memory)

Ubuntu, kernel 2.6.8.1-3-386

Ubuntu, kernel 2.6.8.1-3-386 (recovery mode)

Memory test

Other operating systems:

Windows NT/2000/XP

Figura 3.11 – GRUB é um dos boot managers que acompanham o Linux mais comuns.

Os Windows corporativos (NT, 2000 e XP Professional) possuem um gerenciador de boot que permite um dual boot no computador. Esse programa é chamado NTLDR (NT loader).

Embora existam maneiras, como vimos, de usar mais de um sistema operacional no mesmo computador, é mais comum encontrar computadores com apenas um deles. Vamos estudá-los isoladamente depois de conhecermos seus conceitos básicos.

### 3.9. Sistemas operacionais – conceitos

Sistema Operacional (S.O.) é o programa responsável por manter o computador em funcionamento, responder às requisições do usuário e gerenciar os recursos de hardware da máquina para que trabalhem em “harmonia”.

Todo computador deve ter um sistema operacional para funcionar corretamente. Não é possível um computador funcionar sem sistema operacional. O sistema operacional controla todo o funcionamento do computador.

A maioria das bibliografias especializadas aponta que as funções básicas de um sistema operacional são:

1. Gerenciar os recursos de hardware;
2. Controlar a execução dos programas;
3. Servir de interface entre o usuário e a máquina.

• **Gerenciar os recursos de hardware:** quer dizer que o sistema operacional controla os componentes físicos do computador de forma que a máquina trabalhe corretamente, desde o momento em que o usuário pressiona uma tecla até o aparecimento do referido caractere no monitor.

- **Controlar a execução dos programas:** um software qualquer (como o Word) só é executado com a permissão do sistema operacional e se mantém sob o controle do S.O. até que sua execução termine. Ou seja, enquanto você digita no Word, ele está o tempo todo se reportando ao Windows para acessar memórias, discos, periféricos etc.
- **Servir de interface entre o usuário e a máquina:** é justamente o que o sistema operacional realiza que podemos ver. Tudo o que está à nossa frente na tela, os ícones, as janelas, os comandos etc. Essas são as formas “bonitinhas” de termos acesso aos recursos do computador sem ter de usar a língua dele (binário – 0 e 1). O sistema operacional traduz nossas ações em comandos binários que são entendidos pelo computador e vice-versa, quando a máquina nos dá uma resposta.

### 3.9.1. Componentes do sistema operacional

Um sistema operacional é um programa que gerencia o computador, fazendo-o trabalhar corretamente, gerando ambiente de comunicação entre o usuário e a máquina em si. Essas funções são desempenhadas por dois subsistemas do sistema operacional:

1. Shell;
2. Kernel.

• **Shell:** parte do programa do sistema operacional que cria a interface de comunicação com o usuário. O Shell pode ser gráfico (como o Windows,), quando usa ícones, janelas e um dispositivo apontador (mouse). Nesse caso chamamos de GUI (Interface Gráfica com o Usuário).

O Shell também pode ser textual, quando o usuário conta apenas com o teclado para interagir com o sistema operacional. (Era assim no DOS, e haja comandos para memorizar!)

• **Kernel:** é o núcleo do sistema operacional, eu diria até que Kernel é a “personalidade” do sistema.

O Kernel guarda o funcionamento básico do sistema operacional. É seu componente mais importante. Todo o funcionamento do sistema operacional, desde a forma como se comunica com os dispositivos, até o jeito como armazena seus dados nas memórias e se comunica com o Shell, são definidos em seu Kernel.

### 3.9.2. Tipos de sistemas operacionais

Existem várias classificações, por vários autores, de sistemas operacionais. Resolvi não citar todas neste livro visando não abordar assuntos que são cobrados mais comumente em concursos para analistas de sistemas.

Para nosso entendimento, um sistema operacional pode ser classificado em:

1. Monotarefa ou multitarefa;
2. Monousuário ou multiusuário.

#### 3.9.2.1. Quanto à execução de programas

Um sistema monotarefa não consegue entender as requisições de vários programas ao mesmo tempo. Ele foi desenvolvido para fornecer todos os recursos do computador para apenas um

software (por exemplo, um aplicativo) por vez. O DOS era assim.

Quando um usuário trabalha com um programa qualquer, só poderá utilizar outro programa quando finalizar a utilização do atual; não há possibilidade de execução de dois ou mais programas simultaneamente.

Num sistema multitarefa, vários programas podem usar os recursos do computador ao mesmo tempo (ou quase). Um exemplo de multitarefa são os produtos da família Windows®, que são todos sistemas operacionais multitarefa. Com eles, podemos usar o Word e o Excel em janelas separadas simultaneamente, sem que um interfira nas informações presentes no outro programa. Isso constitui um avanço fantástico.

Imagine um computador executando o Word (o usuário está digitando), mas, por trás disso estão: o CD de música que está tocando, o antivírus verificando a situação, a impressora recebendo dados e imprimindo, a Internet conectada copiando um arquivo para o computador em questão.

Isso só é possível graças ao Windows, que é multitarefa. Posteriormente discutiremos como o sistema operacional manipula as requisições de todos os programas na memória; por ora, apenas esses dois conceitos são necessários.

### 3.9.2.2. Quanto à quantidade de usuários

Um sistema operacional pode ser criado para ser utilizado em um micro doméstico, como os nossos, ou para ser instalado em computadores nas empresas, onde várias pessoas poderão utilizá-lo (inclusive ao mesmo tempo, através de “terminais”).

Quando um sistema operacional é criado para uso pessoal, doméstico, em que, supostamente, será usado por apenas uma pessoa, dizemos que esse sistema é monousuário. Os Windows 95, 98, ME e XP Home são exemplos, e o DOS também era.

Quando um sistema operacional é criado a fim de conseguir controlar a execução de tarefas simultâneas para várias pessoas ao mesmo tempo (isso ocorre, por exemplo, em um computador servidor em empresas), esse sistema é chamado multiusuário.

Nesses sistemas, normalmente há um recurso de identificação do usuário que está usando o computador: cada vez que um usuário liga o computador, o sistema operacional multiusuário lhe solicita uma identificação (o login ou nome de usuário) e exige uma autenticação (a senha) para que este possa acessar os recursos do computador.

Depois da identificação do usuário, o sistema aciona seu Shell para mostrar ao usuário apenas aqueles dados a que ele tem direito (privilegio) de acessar. Sim, os sistemas multiusuário são muito mais seguros que os sistemas monousuário. O Windows 2000, o Linux, o Windows 2008 e o Windows 7 (este, embora para uso doméstico) são exemplos de sistemas multiusuário.

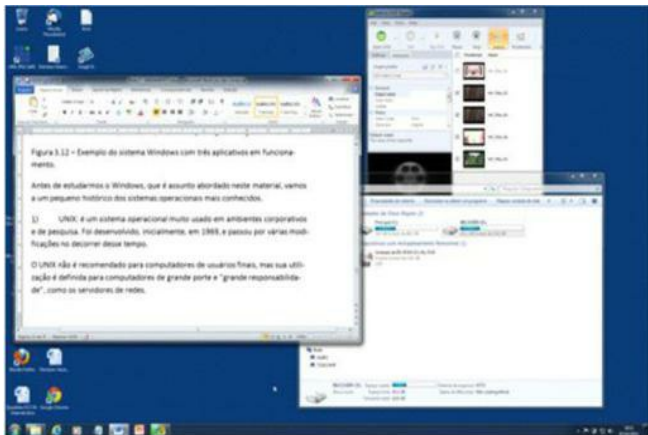
Se um usuário tentar acessar um recurso que não lhe é autorizado, o sistema operacional não permite e avisa ao usuário sobre esse fato. Não é possível burlar a segurança de um sistema multiusuário (pelo menos para nós, meros usuários cotidianos).

Normalmente, há uma pessoa encarregada de gerenciar todos os recursos de um sistema multiusuário que possui uma identificação própria (login e senha) que lhe dá acesso a todos os recursos e possibilidades do sistema. Essa pessoa é chamada **administrador do sistema** e comumente tem uma formação na área de informática.



### 3.9.3. Sistemas operacionais famosos

Sem dúvida alguma, o sistema operacional mais usado hoje em dia é o Windows, desenvolvido pela empresa americana Microsoft, que conta com cerca de 70% do mercado mundial de computadores pessoais. Veja a figura a seguir.



*Figura 3.12 – Exemplo do sistema Windows 7 com três aplicativos em funcionamento.*

Antes de estudarmos o Windows, que é assunto abordado neste material, vamos a um pequeno histórico dos sistemas operacionais mais conhecidos.

- **UNIX:** é um sistema operacional muito usado em ambientes corporativos e de pesquisa. Foi desenvolvido, inicialmente, em 1969, e passou por várias modificações no decorrer desse tempo.

O UNIX não é recomendado para computadores de usuários finais, mas sua utilização é recomendada para computadores de grande porte e “grande responsabilidade”, como os servidores de redes.

O UNIX foi originalmente criado com um Shell de texto, que permitia ao usuário interagir com o sistema sem o uso de um mouse, apenas digitando comandos específicos pelo teclado (uma coisa não muito agradável para a maioria dos usuários, não acha?). Veja a figura a seguir:

```

Greg Wilson@narandine ~/svc/trunk
$ ls
LICENSE.txt      config.mk      docs           ing            license.svc    scraps
Makefile        data          extern        index.svc     press         sites
conf            depend.mk     graphics      lec           print.css     svc.css

Greg Wilson@narandine ~/svc/trunk
$ ls -l
total 64
-rwxr-xr-x 1 Greg Wilson None 1070 Dec 8 15:40 LICENSE.txt
-rwxr-xr-x 1 Greg Wilson None 17030 Jan 19 16:34 Makefile
drwxr-xr-x+ 3 Greg Wilson None 0 Jan 25 13:40 conf
-rwxr-xr-x 1 Greg Wilson None 1807 Feb 1 12:38 config.mk
drwxr-xr-x+ 6 Greg Wilson None 0 Dec 9 16:49 data
-rw-r--r-- 1 Greg Wilson None 8598 Feb 1 12:39 depend.mk
drwxr-xr-x+ 5 Greg Wilson None 0 Jan 3 09:53 docs
drwxr-xr-x+ 3 Greg Wilson None 0 Jan 27 09:38 extern
drwxr-xr-x+ 7 Greg Wilson None 0 Feb 1 13:02 graphics
drwxr-xr-x+ 3 Greg Wilson None 0 Jan 3 09:53 ing
-rwxr-xr-x 1 Greg Wilson None 2912 Jan 3 09:53 index.svc
drwxr-xr-x+ 6 Greg Wilson None 0 Feb 1 12:40 lec
-rwxr-xr-x 1 Greg Wilson None 1302 Dec 8 15:41 license.svc
drwxr-xr-x+ 7 Greg Wilson None 0 Dec 9 16:49 press
-rwxr-xr-x 1 Greg Wilson None 926 Dec 8 15:40 print.css
drwxr-xr-x+ 3 Greg Wilson None 0 Feb 1 08:58 scraps
drwxr-xr-x+ 4 Greg Wilson None 0 Dec 8 15:40 sites
-rwxr-xr-x 1 Greg Wilson None 6950 Jan 16 14:42 svc.css
-rwxr-xr-x 1 Greg Wilson None 1037 Jan 11 13:17 svc.dtd
drwxr-xr-x+ 5 Greg Wilson None 0 Jan 5 10:52 tests
drwxr-xr-x+ 3 Greg Wilson None 0 Feb 1 12:57 tmp
drwxr-xr-x+ 3 Greg Wilson None 0 Feb 1 12:57 util
drwxr-xr-x+ 4 Greg Wilson None 0 Feb 1 13:00 web

Greg Wilson@narandine ~/svc/trunk
$ =

```

Figura 3.13 – Tela do UNIX (interface textual).

O UNIX serviu como base para a criação de diversos outros sistemas operacionais de empresas diferentes, devido a seu projeto arrojado e sério. LINUX, SOLARIS, FREEBSD e MacOS (dos computadores da Apple) são alguns dos sistemas operacionais criados com base no UNIX.

O UNIX é um sistema operacional multiusuário, multitarefa e de servidor. Não é muito comum encontrá-lo em computadores pessoais como os que usamos, mas em computadores de grande porte.

- **DOS:** Sistema operacional criado pela Microsoft em meados de 1980. Foi o grande estopim para alavancar as vendas dos computadores pessoais como conhecemos hoje. Veja a figura a seguir (e mate as saudades).

```

C:\Users\Joao Antonio>Dir
O volume na unidade C é Principal
O Número de Série do Volume é 94DB-4E97

Pasta de C:\Users\Joao Antonio

02/10/2012  15:32    <DIR>          .
02/10/2012  15:32    <DIR>          ..
13/07/2012  09:23    <DIR>          Contacts
02/10/2012  16:11    <DIR>          Desktop
15/03/2012  21:19    <DIR>          Diablo-III-8370-ptBR-Installer
23/09/2012  08:42    <DIR>          Documents
21/09/2012  12:47    <DIR>          Downloads
13/07/2012  09:23    <DIR>          Favorites
13/07/2012  09:23    <DIR>          Links
13/07/2012  09:23    <DIR>          Music
20/09/2012  18:12    <DIR>          Pictures
13/07/2012  09:23    <DIR>          Saved Games
13/07/2012  09:23    <DIR>          Searches
13/07/2012  09:23    <DIR>          Videos
            0 arquivo(s)                0 bytes
            14 pasta(s)        333.084.160.000 bytes disponíveis

C:\Users\Joao Antonio>

```

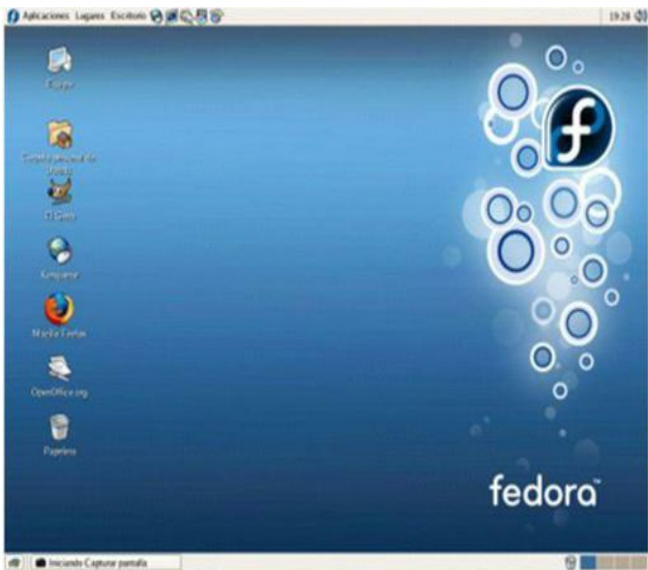
*Figura 3.14 – Sentiu falta dele nos últimos concursos?*

Com poucos recursos e uma interface desagradável, o DOS era a única opção para os usuários de PCs até meados de 1990. O DOS era monotarefa, monousuário e foi desenvolvido para ser usado apenas em computadores pessoais.

### 3.9.4. Linux – O “Patinho” Feio?

O Linux é um sistema operacional desenvolvido com base no UNIX. Foi criado inicialmente (em 1991) pelo estudante finlandês Linus Torvalds e hoje é uma sensação no mundo da informática. O Linux foi desenvolvido para dar suporte a qualquer tipo de computador, não importa se é apenas um computador pessoal ou um servidor no centro nervoso de uma empresa multinacional.

Conheça o Linux:



*Figura 3.15 – KDE (uma das interfaces gráficas mais usadas atualmente no Linux).*

Apesar de ter sido baseado inicialmente em interfaces de texto, como o UNIX original, o Linux é hoje apresentado com várias interfaces gráficas. Podemos escolher qual será a cara do Linux que utilizaremos.

O Linux é um software livre: isso quer dizer que ele respeita (e, a rigor, sempre respeitará, em suas futuras versões) as quatro “liberdades” dos usuários, permitindo que eles USEM o software, COPIEM-NO e distribuam-no para quantos usuários quiserem, ESTUDEM-NO e até MODIFIQUEM-NO, criando novos e melhorados softwares com base nele.

Claro que, para ter acesso a estudar e modificar o Linux, o seu código-fonte (a receita que descreve como ele foi programado, ou seja, como foi feito) deve ser liberado (acessível) para todos os usuários!

Mais informações sobre o Linux podem ser conseguidas numa apostila completa que escrevi

sobre o programa. Acesse [www.euvoupassar.com.br](http://www.euvoupassar.com.br) e pegue a apostila na seção Materiais Avulsos (na Home do site).

### 3.9.5. iOS – Sistema operacional do iPhone e iPad

A Apple®, fabricante dos computadores Macintosh e de dispositivos famosos como o Smartphone iPhone e o Tablet iPad, também é fabricante de sistemas operacionais para estes dispositivos.

O sistema operacional do iPhone e iPad é chamado de iOS.



*Figura 3.16 – Tela do iOS num iPad.*

Com ele, é possível acessar a Internet, instalar programas, jogar, e, claro... fazer ligações telefônicas (no iPhone). É um sistema multitarefa (permite a execução de vários programas ao

mesmo tempo) e é muito fácil de aprender.

### 3.9.6. Android – o sistema operacional da Google

A empresa Google®, que, entre outras coisas, é dona do site [www.google.com](http://www.google.com), o maior site de busca da Internet, também “deu as caras” nos sistemas operacionais!

**Android** é o nome dado ao sistema operacional que o Google criou para dispositivos portáteis, como smartphones e tablets. Ele é o principal concorrente do iOS.

O Android traz uma série de recursos interessantes, que rivalizam com o iOS, e tem a vantagem de ser um software de código aberto (open-source), ou seja, seu código-fonte é distribuído, permitindo que qualquer programador possa estudar e modificar o Android, gerando cópias “melhoradas” do sistema.

Ainda há uma discussão mais acalorada, especialmente para os mais “xiitas” no assunto se o Android pode ser considerado um software livre (um programa que respeita as quatro liberdades do usuário: usar, copiar, estudar, modificar) em todo o seu conteúdo (porque há partes do Android que não foram “liberadas”)... Mas isso é para outro momento!

Então, por segurança, considera-se o Android um sistema Open Source (código aberto)! Ahhh, quase ia me esquecendo: o Android é um derivado do Linux!

Sim! Sim! O sistema Android é composto pelo núcleo (e algumas outras partes) do sistema Linux! Então não seria errado dizer que o Android é um “tipo” de Linux! E ele foi criado especificamente para dispositivos móveis (smartphones e tablets).



*Figura 3.17 – Android 4.0 sendo usado no Tablet Samsung Galaxy 10.1®.*

Bem, com isso conhecemos os principais sistemas operacionais da atualidade, mas ainda vamos dar ênfase ao mais usado e, por isso, mais cobrado em concursos públicos: o Windows.

#### **4.1. Pequeno histórico do Windows**

O sistema operacional Windows foi desenvolvido (e ainda é) pela Microsoft®, que começou seu projeto no final da década de 1980.

Inicialmente, a Microsoft criou o Windows apenas como uma “interface” para o sistema operacional daquela época (o DOS); portanto, as primeiras versões do Windows não eram sistemas operacionais, eram classificados apenas como ambientes gráficos (eles eram apenas uma “cara” bonita para o DOS).

Alguns exemplos de ambientes gráficos que rodavam (eram executados) no sistema operacional DOS: Windows 3.0, Windows 3.1 e Windows 3.11.

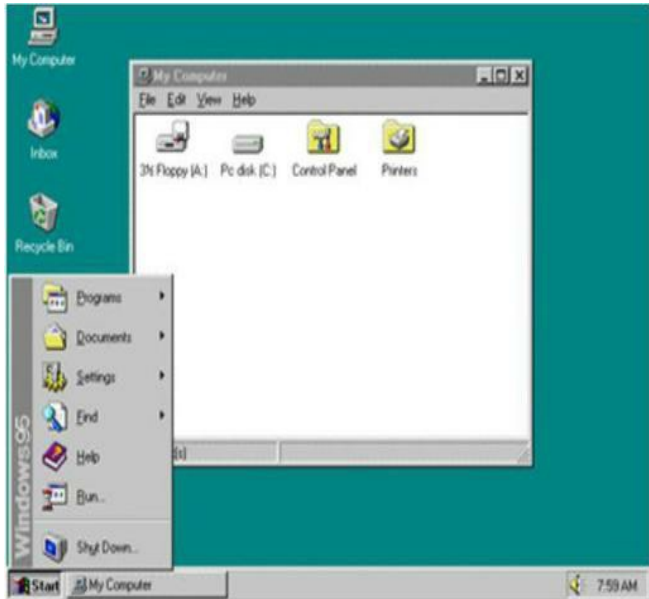




*Figura 4.1 – Windows 3.1 (usado em conjunto com o DOS para controlar a máquina).*

Em 1995, a Microsoft lançou o Windows 95, seu primeiro sistema operacional gráfico, ou seja, cuja interface não era baseada em texto, como o DOS sempre foi. O Windows 95 era autônomo, pois não dependia do DOS.

Lembre-se! Qualquer questão (se houver mais alguma sobre esse assunto daqui para frente) que fale sobre o Windows “precisar” do DOS está errada! A partir do Windows 95, todos são considerados sistemas operacionais.



*Figura 4.2 – Windows 95 – mudança radical.*

A principal mudança entre o DOS e o Windows 95 é que esse novo sistema foi compilado (projetado) para rodar em 32 bits, diferentemente do DOS, que era um sistema operacional de 16 bits. Isso significa que suas instruções (comandos) são dadas à CPU em blocos de 32 em 32 bits.

Naquela época, o hardware já suportava tal capacidade, mas não havia na Microsoft nenhum software que pudesse usufruir de todo esse “poder”.

Vale salientar que, hoje em dia, como vimos no capítulo de hardware, isso é passado, porque as CPUs (processadores) da atualidade são todas de 64 bits.

Vejamos uma pequena lista com as características principais de algumas das versões anteriores do sistema Windows:

- **Windows 95:** lançado em 1995 para o mercado doméstico, este sistema operacional usava o sistema de arquivos FAT16 para gerenciar as partições dos discos rígidos. Iniciou a era Plug and Play (detecção automática de dispositivos para facilitar a instalação deles) – é “intitulado” o primeiro sistema operacional com suporte à tecnologia Plug and Play da história! (Que fantástico!)

Houve uma versão posterior desse sistema, chamada Windows 95 OSR 2 (também chamado de Windows 95 B), que já trazia suporte ao sistema de arquivos FAT32.

- **Windows 98:** lançado em 1998 (claro), este sistema operacional já utilizava o sistema de arquivos FAT32. Surgiu apenas como uma evolução natural do Windows 95. Possuía um programa conversor de FAT16 para FAT32.

Foi a partir dessa versão que o Internet Explorer (programa navegador) e o Outlook Express (programa de correio eletrônico) passaram a fazer parte do próprio sistema operacional Windows. Ou seja, o Windows 98 apresentava maior interação com a Internet que o seu antecessor.

- **Windows ME (Edição do Milênio):** foi lançado em 2000 e não trouxe muitas novidades em relação ao Windows 98. Pelo menos teve lugar de destaque: foi considerado o pior de todos os Windows (o mais problemático deles).

- **Windows XP:** foi lançado em 2001 e permanece, ainda hoje, sendo o mais usado dentre todas as versões do Windows (e ainda é o mais cobrado em provas!).

Esta edição do livro não traz mais o Windows XP como versão principal, mas se você vai enfrentar algum concurso que exija esta versão do sistema, poderá pegar a apostila sobre Windows XP gratuitamente no site do [www.euvoupassar.com.br](http://www.euvoupassar.com.br), na seção de Materiais Avulsos (que aparece logo na home page!)

- **Windows Vista:** sucessor do Windows XP, esta versão do Windows foi lançada em 2006 e sinceramente merece ser esquecida! Muito “pesada” e cheia de problemas que geraram inúmeras reclamações, o Windows Vista foi logo substituído pelo Windows 7!

- **Windows 7:** em outubro de 2009, a Microsoft lança sua novíssima versão. É esta a versão atual do Windows e, claro, é a que abordaremos neste capítulo!

Mas a Microsoft investiu também no mercado corporativo de computadores, ou seja, nas empresas. Os sistemas operacionais que ela criou para essa finalidade, no decorrer da história, foram:

- **Windows NT:** criado em 1992 e aperfeiçoado em 1995, o Windows NT era um sistema operacional com diversos recursos para a utilização em ambiente corporativo. É um sistema multiusuário, que suporta os sistemas de arquivos NTFS e FAT16 (o Windows NT não entendia FAT32).

O Windows NT não tinha semelhanças internas com o Windows 95. Eram dois produtos completamente diferentes, originados de projetos muito distintos.

- **Windows 2000:** sucessor natural do Windows NT, este sistema operacional trouxe as vantagens de entender o sistema de arquivos FAT32 (além de continuar entendendo o FAT16) e já utilizar o sistema de arquivos NTFS 5 (atualização do NTFS do NT).

- **Windows 2003 Server:** sucessor do Windows 2000 Server, usado em servidores robustos. Este sistema já apresenta versões em 64 bits (para os processadores mais recentes).

- **Windows 2008 Server:** sucessor do Windows 2003 Server. É a versão atual do sistema operacional Windows para servidores.

## 4.2. Características básicas do sistema Windows

O sistema operacional Windows (não importando a versão exatamente) tem uma série de características que devem ser apresentadas ao concursando e não podem ser esquecidas na hora de fazer a prova:

O Windows é um sistema operacional gráfico: isso significa que sua interface (ou seja, sua “cara”) é baseada em itens visuais, como ícones, janelas, menus. Não é necessário que o usuário digite comandos como os comandos usados no DOS e UNIX para acionar o sistema. É só usar os itens que se apresentam de forma “bonitinha” na tela.

O Windows sempre foi (historicamente) um sistema operacional de 32 bits: isso significa que ele foi criado para controlar máquinas com processadores que usam essa tecnologia (32 bits no barramento de dados).

É bom lembrar que já há algumas versões do sistema Windows, como o Windows 2003 Server e, mais recentemente, o Windows Vista e o Windows 7, com exemplares compilados para 64 bits. Esses sistemas operacionais são otimizados para serem usados em CPUs de 64 bits.

O Windows usa multitarefa preemptiva: isso quer dizer que o Windows permite a execução de várias tarefas ao mesmo tempo (pelo menos, faz aparentar isso para o usuário). A multitarefa preemptiva é um sistema que permite que várias janelas de vários programas sejam apresentadas ao usuário, como se todos estivessem sendo “executados” ao mesmo tempo.

Na verdade, o que acontece é que o Windows fica “chaveando” a execução de tarefas na CPU de forma bem rápida (isso porque, só há uma CPU no micro), fazendo parecer que pode fazer tudo ao mesmo tempo. Ele fica mais ou menos como um guarda de trânsito, fazendo: “Impressora, é sua vez...”, “Pare!”, “Agora é a vez do Word, pronto, pode passar”, “Agora é o Excel que vai usar a CPU! Prooonto... Deixe de ser egoísta”, “Pare”, “Agora é a vez do Word de novo...” e assim por diante.

Em resumo, na multitarefa preemptiva, é o sistema operacional que controla de quanto tempo (e de quantos recursos) um programa pode dispor um determinando momento.

O Windows suporta Plug and Play: significa que a instalação de equipamentos Plug and Play pode ser realizada de forma simples no Windows, que entende perfeitamente esse sistema.

Os sistemas Windows são dotados da tecnologia Plug and Play, que permite que eles reconheçam automaticamente equipamentos de hardware no momento de sua instalação, facilitando muito a vida dos usuários na hora de adicionar um novo equipamento ao computador.

Lembre-se (já vimos isso): Plug and Play é uma “filosofia” desenvolvida em conjunto com vários fabricantes de hardware e software para que um computador consiga reconhecer automaticamente um equipamento que foi instalado fisicamente nele (por exemplo, uma nova impressora).

Funciona assim: uma impressora Plug and Play (todas, hoje em dia) possui um chip de memória ROM com suas informações básicas de identificação, o sistema operacional Windows simplesmente “lê” esse chip para reconhecer a impressora. Já vimos anteriormente como instalar um hardware no ambiente Windows com o uso da tecnologia Plug and Play.

#### 4.2.1. Como o Windows entende as unidades

Uma das principais “responsabilidades” de um sistema operacional é, sem dúvida, o gerenciamento de arquivos (dados armazenados em memórias permanentes, como vimos anteriormente). Um sistema operacional tem de ser capaz de permitir ao usuário realizar diversas ações com arquivos, pastas e unidades de armazenamento (como copiar, formatar, excluir etc.).

Com relação às unidades de armazenamento, ou simplesmente *unidades*, que são as memórias permanentes em nosso computador, como já foi visto, o Windows atribui um identificador único a cada uma delas, baseado em uma nomenclatura própria. Cada unidade recebe *uma letra seguida do sinal de dois pontos (:)*. Cada unidade instalada no computador receberá uma letra diferente.

## ➤ Unidades de Disco Rígido (2)



Principal (C:)



RECOVERY (F:)

## ➤ Dispositivos com Armazenamento Removível (2)



Unidade de  
BD-ROM (D:) My  
DVD



Dispositivo de  
Armazenamento  
Digital Seguro  
(E:)

*Figura 4.3 – Ícones das unidades e suas respectivas identificações.*

As unidades **A:** e **B:** sempre serão destinadas a dispositivos de disquete (até o presente momento, pelo menos) e é justamente por isso que não aparece nenhuma unidade com essas letras no Meu Computador mostrado acima! Não existem mais unidades de disquete!

A unidade denominada **C:** está reservada para uma partição de Disco Rígido (HD) – mais precisamente, a primeira partição do primeiro disco rígido. As demais letras das unidades serão destinadas aos outros equipamentos que serão instalados no computador (ou demais partições do disco rígido).

É justamente nas unidades que estão os arquivos e as pastas do seu computador.

Alguns computadores apresentarão mais unidades, outros apresentarão menos unidades (isso

dependerá, exclusivamente, do número de equipamentos de memória auxiliar que foram instalados em seu computador).

Quanto ao exemplo mostrado na figura anterior, temos duas *unidades de Disco Rígido (HD)*, sendo *C:* e *F:*, além de uma *unidade de BD/DVD/CD* (chamada de *D:*) e finalmente uma *unidade de Cartão de Memória do tipo SD* (intitulada *E:*).

#### 4.2.2. Como o Windows trata os arquivos

Continuando a forma como o Windows gerencia os dados armazenados em unidades de disco (dados que são conhecidos como arquivos), segue uma explicação básica de como os próprios arquivos são entendidos pelo sistema operacional.

Um arquivo pode ser classificado como arquivo de dados (que contém dados normalmente feitos pelo usuário) ou arquivo de programa (que contém instruções a serem executadas pelo sistema operacional). Os arquivos do Word e do Excel, como os que criamos cotidianamente, são arquivos de dados, mas os próprios Word e Excel são armazenados em arquivos de programas (chamados de arquivos executáveis).

Há algumas regras que devem ser seguidas para nomear (e renomear) um arquivo ou uma pasta no sistema operacional Windows. Aqui vão elas:

1. Um nome de arquivo ou pasta deve ter até 255 caracteres.
2. Não podem ser usados os seguintes caracteres: \* (asterisco), “ (aspas), > (sinal de maior), < (sinal de menor), : (dois pontos), / (barra), | (barra vertical), \ (barra invertida) e ? (interrogação).
3. Não pode haver dois objetos com o mesmo nome no mesmo diretório (pasta).
4. Arquivos possuem extensão (chamo, carinhosamente, de “sobrenome”), que é um conjunto de três caracteres (normalmente) e serve para identificar o tipo de um arquivo. Isso não é uma “convenção”.

Quem atribui a extensão ao arquivo é o próprio programa que o cria, como o Word e o Excel, por exemplo. Normalmente, no Windows, as extensões estão ocultas para o usuário, mas é possível solicitar ao programa que as mostre. Verifique a seguir alguns arquivos com extensões diversas.



*Figura 4.4 – Alguns arquivos e suas extensões.*

#### 4.2.3. Extensões dos tipos de arquivos mais comuns no Windows

Aqui estão as extensões de alguns dos principais tipos de arquivos conhecidos de quem usa o Windows. Resolvi apresentar mais que uma simples lista delas: junto com cada extensão, há uma explicação mais aprofundada do tipo de arquivo correspondente.

##### 4.2.3.1. Arquivos usados no dia a dia

##### Extensão DOC (Documento do Microsoft Word)

São os arquivos criados pelo programa Microsoft Word. Esses arquivos podem conter diversos tipos de dados, como textos, figuras, tabelas, efeitos de bordas etc.

O Word é o programa de texto que vem no conjunto Microsoft Office (conjunto de programas de escritório da Microsoft).

Os documentos do Word são arquivos que podem possuir algum conteúdo interno de programação (as chamadas macros). As macros são criadas em uma linguagem de programação que acompanha todos os programas do pacote Microsoft Office (do qual o Word faz parte): VBA (Visual Basic for Applications).



As macros são o principal motivo de os arquivos do Word serem considerados extremamente vulneráveis a infecção por vírus de computador. Os vírus mais comuns da atualidade são conhecidos como vírus de macro e podem infectar os documentos do Word com facilidade, por serem escritos na mesma linguagem de programação que as macros do programa, VBA.

### **Extensão DOCX (Documento do Microsoft Word 2007)**

A penúltima versão do Microsoft Word (lançado em janeiro de 2007) trouxe inúmeras novidades, como veremos no capítulo específico sobre o programa. Uma de suas principais diferenças em relação às versões anteriores do programa é a mudança na extensão do arquivo (que, claro, vem acompanhada da mudança da estrutura interna do mesmo).

Um documento criado pelo Word 2007 não é suportado por versões anteriores do programa. As alterações feitas no interior dos arquivos DOCX resultam em arquivos completamente incompatíveis com as versões anteriores do Word.

O Word 2010 (última versão deste programa e a versão que é abordada neste livro) também utiliza extensão DOCX. Tanto o Word 2010 quanto o Word 2007 conseguem salvar e abrir arquivos no formato DOC.

### **Extensão ODT (Documento do LibreOffice Writer)**

O principal concorrente (em programa de texto) do Word é o *Writer*. O Writer faz parte de um conjunto de programas de escritório chamado *LibreOffice* (antigamente conhecido como *BrOffice*), uma variação do mundialmente conhecido OpenOffice.

O formato ODT traz basicamente todos os recursos do DOC e DOCX. A maioria dos documentos que podemos escrever em DOC/DOCX (como apostilas, livros, documentos técnicos, trabalhos acadêmicos etc.), pode ser feita no LibreOffice também!

A principal vantagem do LibreOffice em relação ao Microsoft Office, do qual o Word faz parte, é que aquele é um software livre!

Ahhh! O formato ODT também é muito menos suscetível a vírus de macro que o formato DOC/DOCX, porque o LibreOffice não usa VBA em seu interior!

### **Extensão XLS (Pasta de Trabalho do Microsoft Excel)**

Os arquivos que possuem a extensão XLS são criados pelo programa Microsoft Excel, que também faz parte do Microsoft Office, o pacote de programas de escritório da Microsoft.

Ao contrário do que muita gente pensa, a descrição desse tipo de arquivo é realmente “Pasta de Trabalho do Microsoft Excel” e não “Planilha do Microsoft Excel”! Isso é um erro comum no dia a dia, mas que não é cometido (nem perdoado) pelas bancas examinadoras de provas de concursos públicos!

Atenção: especialmente se a banca examinadora for a FCC (Fundação Carlos Chagas), é necessário ter especial atenção a esse fato!

Planilha é o nome dado a cada uma das “folhas” que possuem colunas e linhas e que serão preenchidas por valores pelo usuário. Ao conjunto de todas as planilhas (e pode haver muitas em um arquivo do Excel), dá-se o nome de *Pasta de Trabalho* (ou seja, o arquivo em si).

De forma semelhante aos arquivos DOC, os arquivos XLS admitem macros em VBA, o que também os torna bastante vulneráveis aos vírus de macro que se propagam pela Internet. Aliás,

vulnerabilidade aos vírus é característica de todos os tipos de arquivos do Microsoft Office.

### **Extensão XLSX (Pasta de Trabalho do Microsoft Excel 2007)**

A Microsoft também alterou o formato interno do arquivo do Excel no Office 2007 e 2010. O arquivo XLSX não consegue ser aberto por versões anteriores do Microsoft Excel.

Dentre as novidades, há um número maior de linhas e colunas para cada planilha. Recursos de formatação mais avançados também foram adicionados ao programa. Conheceremos mais do Excel 2010 num capítulo próprio.

### **Extensão ODS (Pasta de Trabalho do LibreOffice Calc)**

O LibreOffice Calc é o principal concorrente do Microsoft Excel. Ele também serve para construir planilhas de cálculos complexas.

O Calc possui quase todos os recursos do Excel (que está no “mercado” há mais tempo) e também é um software livre (aliás, como todo o conjunto LibreOffice).

Os arquivos salvos pelo Calc apresentam a extensão ODS.

### **Extensão PPT (Apresentação de Slides do Microsoft PowerPoint)**

Arquivos com extensão PPT pertencem ao Microsoft PowerPoint, o programa para criação e edição de apresentação de slides que faz parte do pacote Microsoft Office.

Os arquivos PPT podem ser alterados por completo. Ou seja, quem possui um arquivo PPT, tem acesso a todo o seu conteúdo e efeitos, podendo alterá-lo completamente em vários aspectos.

### **Extensão PPS (Apresentação de Slides do Microsoft PowerPoint)**

Esse arquivo é normalmente usado para enviar apresentações de slides por e-mail. Esse arquivo contém a apresentação inteira, mas essa apresentação sempre será aberta no modo Apresentação de slides, e não no modo Normal, como o formato PPT.

Em poucas palavras, quando um arquivo PPS é aberto, ele já é aberto apresentando-se, sem dar acesso ao corpo da apresentação em si, que permite a alteração de seu conteúdo.

Quando recebemos um e-mail com um arquivo PPS em anexo, ao abri-lo, ele sempre será imediatamente executado! Um arquivo no formato PPT abriria o PowerPoint primeiro para que, dentro dele, o usuário pudesse alterá-lo ou executá-lo em forma de apresentação.

Como esses arquivos são apresentações, eles também podem transferir vírus de macro. Cuidado!

### **Extensão PPTX (Apresentação de Slides do Microsoft PowerPoint 2007)**

Outro formato novo, incompatível com as versões anteriores do Microsoft PowerPoint. Um arquivo PPTX traz novidades em vários quesitos, como desenhos, animações e muito mais.

É utilizado pelas versões 2007 e 2010 (as mais recentes) do PowerPoint.

### **Extensão TXT (Arquivo de Texto Puro)**

Os arquivos de texto simples (ou texto puro), representados pela extensão TXT, são compostos, como o próprio nome diz, apenas de texto simples (letras e números), ou seja, caracteres simples.

Esses arquivos são bem pequenos (em bytes) e não aceitam nenhum tipo de formatação

(negrito, itálico, sublinhado, cor da fonte, parágrafo, entre outros). Não é possível, também, nesse tipo de arquivo, colocar qualquer tipo de objeto que não seja texto, como fotos, sons, imagens, efeitos especiais etc.

Os arquivos com extensão TXT são arquivos criados e editados por programas editores de texto, como o Bloco de Notas no Windows.

Arquivos com extensão TXT não podem carregar vírus de computador. Seu conteúdo é muito simplificado e os vírus não têm, nele, um habitat adequado para existir. (Você verá, mais adiante, que os vírus precisam de arquivos com algum conteúdo executável, como as macros nos documentos do Office.)

### **Extensão EXE (Arquivo Executável)**

Os arquivos que apresentam a extensão EXE são chamados de arquivos executáveis e, em poucas palavras, são programas. Os arquivos executáveis são compostos por diversas instruções (comandos) escritas no que se chama linguagem de máquina, ou seja, na linguagem que o computador entende e obedece.

Todos os programas que utilizamos (Word, Excel, Calculadora, Internet Explorer, Paciência, Campo Minado etc.) são armazenados na forma de arquivos executáveis, ou seja, todos os programas de nosso computador são, na verdade, arquivos com extensão EXE.

Os arquivos EXE são os programas prontos para serem usados, escritos diretamente na linguagem que o computador compreende. Mas antes de se tornar arquivo executável, esse programa foi escrito por alguém de carne e osso, em uma linguagem de programação não tão estranha, como Pascal, C++, Java etc. Esse arquivo, ainda em sua forma legível para os humanos, é chamado arquivo-fonte porque dá origem ao arquivo executável em si.

Os arquivos-fonte têm diversos tipos de extensão, de acordo com o programa em que foram escritos, mas, em sua maioria, são apenas arquivos de texto com outra extensão. Já os arquivos executáveis estão escritos em uma linguagem completamente estranha para nós, e perfeitamente clara e compreensível para o computador.

Todos os arquivos EXE são vetores de transmissão de vírus de computador. É justamente nos arquivos EXE que a maioria dos vírus de computador encontra seu habitat perfeito. Se você receber, por e-mail ou por outro meio, um arquivo com extensão EXE, desconfie. Ou melhor, apague-o imediatamente!

### **Extensão PDF (Arquivo do Adobe Acrobat)**

Os arquivos PDF são muito comuns para disponibilizar informações de texto na Internet como apostilas, material de consulta técnica, editais, documentos oficiais, monografias, teses etc. O uso exagerado desse tipo de arquivo tem uma razão: ele foi criado para não ser alterado.

Os arquivos com extensão PDF foram criados para serem usados em cartórios, pois, como a ideia diz, depois de digitalizados (transformados em arquivos) os documentos não deveriam conceber mais alterações.

Para criar um arquivo PDF, deve-se possuir um programa com essa capacidade, como o Adobe Acrobat (desenvolvido e vendido pela empresa Adobe – é isso mesmo, ele não é gratuito). Para ler os arquivos PDF, porém, usa-se um programa gratuito chamado Adobe Reader, ou Adobe Acrobat Reader, que realiza apenas a operação de leitura desses arquivos.

Dentre as vantagens do formato PDF, podemos citar:

1. O arquivo PDF não pode ser alterado (isso é interessante para a segurança e a autenticidade do documento – quem gosta dessa característica são os autores de documentos na Internet, como eu).
2. O arquivo PDF mantém as mesmas configurações visuais do arquivo que o originou (tamanhos e tipos de fonte, figuras, margens, tamanho da página etc.), independentemente de a pessoa que vai ler o arquivo possuir o programa que criou originalmente aquele arquivo ou não.
3. O arquivo PDF é normalmente menor (em bytes) que o arquivo que o origina (ou seja, o algoritmo usado para criar o PDF também compacta seus dados).

Apesar de ter sido criado para não ser alterado, dizer que a segurança de um arquivo PDF é intransponível é exagero. Já há diversas formas de fazer a “engenharia reversa” em um arquivo PDF, ou seja, transformá-lo de volta no texto que o originou (arquivo do Word, Excel, PowerPoint etc.).

PDF significa **Portable Document File** (Arquivo de Documento Portátil).

### **Extensão RTF (Documento Genérico)**

Os arquivos com extensão RTF são documentos de texto que admitem formatação (negrito, itálico, sublinhado, tipo e tamanho de fonte etc.) e recursos especiais como figuras e tabelas (como os arquivos DOC do Word). Os arquivos RTF são legíveis por qualquer processador de texto atual.

Esse formato de arquivo era muito usado para a troca de documentos entre empresas ou usuários com diferentes programas processadores de texto.

Atualmente, não se vê muita utilidade no formato RTF simplesmente porque o formato DOC já se popularizou. Todos os programas processadores de texto do mercado basicamente conseguem abrir e salvar arquivos no formato DOC, tornando-o hegemônico.

Na época, porém, em que o Word não estava tão “em evidência”, o formato RTF garantia a troca de arquivos de texto entre programas diferentes, como o Microsoft Word, o Corel WordPerfect, o Microsoft Works, entre outros.

### **Extensão LNK (Atalho)**

A extensão LNK é menos comum de se ver. Pois, mesmo configurando o Windows para exibir as extensões dos arquivos, essa fica oculta. Arquivos LNK são atalhos. Sim, atalhos! Aqueles ícones com o quadradinho branco e a setinha curva no canto inferior esquerdo.

Atalhos são apenas arquivos que apontam para outros arquivos. Ou melhor, são arquivos que apontam para qualquer tipo de objeto selecionável no Windows, como arquivos, pastas, impressoras, unidades de disco, páginas da Internet etc.

### **Extensão ZIP (Arquivo Compactado)**

Os arquivos que possuem a extensão ZIP são arquivos cujo conteúdo foi compactado por programas como o **WinZip**. Arquivos ZIP podem conter diversos outros arquivos em seu interior e, normalmente, apresentam seu tamanho (em bytes) menor que o da soma dos arquivos neles contidos.

Em suma, se você possui muitos arquivos de diversos tipos que, somados, apresentam 300 KB de tamanho, quando eles forem compactados serão transformados em apenas um arquivo com extensão ZIP de tamanho, 200 KB, por exemplo. Veremos mais sobre o funcionamento dos compactadores no próximo capítulo do livro.

Os arquivos ZIP podem trazer vírus aos nossos computadores? Claro que sim! Se um arquivo ZIP for composto de vários outros arquivos, basta que um desses arquivos (no interior do ZIP) esteja infectado e seja aberto para desencadear a infecção em nossos computadores.

### **Extensão RAR (Arquivo Compactado)**

Existem outros programas de compactação/descompactação além do WinZip e seus equivalentes: o WinRAR é um deles.

O WinRAR usa um algoritmo (conjunto de processos matemáticos e lógicos) diferente do algoritmo usado nos arquivos ZIP, o que os torna diferentes internamente. (Isso resulta em o programa WinZip não ser capaz de abrir/manipular arquivos RAR, embora o programa WinRAR tenha capacidade de compactar e descompactar arquivos em ZIP naturalmente.)

Como a ideia dos dois é a mesma, os arquivos RAR podem trazer, em seu interior, algum arquivo infectado por vírus, portanto cuidado!

#### **4.2.3.2. Arquivos de multimídia**

Os arquivos a seguir são compostos apenas de dados, e não de instruções, o que os torna habitats inadequados para vírus de computador. Ou seja, nenhum dos tipos de arquivo a seguir tem potencialidades de infecção por vírus.

É possível que, em um futuro próximo, novos vírus de computador venham a ser criados para afetar tais arquivos, mas até agora, eu me arrisco a dizer que esses arquivos, com exceção de alguns que irei citar, são seguros.

### **Extensão BMP (Imagem de Bitmap)**

Os arquivos com extensão BMP são imagens de bitmap (mapa de bits). Esses arquivos são figuras compostas de pequenos quadradinhos (os pixels) que podem apresentar milhões de cores, o que os torna ideais para armazenar imagens fotográficas.

Esses arquivos podem ser abertos por qualquer tipo de programa para editar fotografia (como o Adobe Photoshop e o Adobe Fireworks), mas são famosos por ser o formato de arquivo usado no programinha *Paint* que acompanha o Windows.

Os arquivos BMP não são compactados, ou seja, cada pixel (quadrado colorido) é armazenado completamente no arquivo, ou seja, cada um dos pontos que formam a fotografia vai gastar um número de bytes próprio (e fixo).

### **Extensão JPG (ou JPEG) (Imagem de Bitmap Compactada)**

Esses arquivos são, na verdade, arquivos BMP que passaram por um processo de compactação em sua estrutura de dados. Os arquivos JPEG também podem apresentar milhões de cores.

Todas as fotos JPG são baseadas em arquivos BMP e sua estrutura de armazenamento. A grande maioria das máquinas fotográficas digitais já armazena as fotos tiradas em formato JPG. As figuras em formato JPG são muito comuns em páginas da Internet.

Os arquivos JPG são bem menores em tamanho em relação aos arquivos em formato BMP, ou seja, uma foto em formato BMP pode ser 10 a 20 vezes maior, em bytes, que sua equivalente em formato JPG. Essa compactação acontece com perdas, ou seja, os arquivos JPG deixam de armazenar alguns dados da fotografia que são armazenados no arquivo BMP.

Na verdade, JPEG é o nome da empresa (equipe) que criou e atualiza o processo matemático (algoritmo) que compacta (diminui o tamanho) os arquivos de imagem BMP. JPEG significa **Joint Picture Experts Group** (Grupo de Especialistas em Imagens Juntas) – é, eu sei que a sigla não ajuda!

### **Extensão GIF (Imagem de Bitmap Compactada)**

Os arquivos GIF também são imagens de bitmap usadas para armazenar desenhos criados em softwares de desenhos. Os arquivos com extensão GIF são compactados e apresentam tamanhos (em bytes) menores que os equivalentes em BMP.

Embora tenham características semelhantes aos arquivos JPG, os arquivos GIF apresentam menos cores (apenas 256), tornando-os, por causa disso, não recomendáveis para fotografias. Normalmente, nas páginas da Internet, os arquivos GIF são usados para representar figuras pequenas, como papéis de fundo, logomarcas e detalhes das páginas.

As figuras GIF podem ser animadas, recurso muito utilizado nas primeiras páginas da Internet, no final da década de 1990.

### **Extensão PNG (Imagem de Bitmap)**

Arquivos com extensão PNG são fotografias, porém mais complexas que as imagens em BMP ou em JPEG. Um arquivo PNG admite uma série de recursos avançados que não são suportados pelos formatos JPEG e GIF; portanto, esse tipo especial de figura é cotado para ser o substituto dos dois anteriores nas páginas da Internet em alguns anos.

PNG significa **Portable Network Graphics** (Gráficos Portáteis para Redes), e isso indica muito bem a intenção desse formato de arquivo em estar intimamente ligado com a Internet (redes).

A maioria dos programas de edição de fotografias é capaz de abrir e manipular arquivos com a extensão PNG, que será mais conhecido daqui a algum tempo. A única vantagem quanto ao uso do JPG em detrimento do PNG é o seu tamanho. (Arquivos JPG ainda são menores!)

### **Extensão AVI (Vídeo)**

Os arquivos AVI são **vídeos simples**. Podem conter imagem em movimento e som.

Os arquivos AVI não são compactados, o que os torna muito grandes (em bytes) – ou seja, os arquivos AVI conseguem armazenar, sem perdas, todos os dados do vídeo gravado, sem exceção!

Os arquivos no formato AVI podem ser abertos em qualquer programa de vídeo, incluindo aqueles que vêm junto com o Windows (Windows Media Player).

### **Extensão MPG (ou MPEG) (Vídeo Compactado)**

O JPG está para o BMP assim como o MPG está para o AVI. Os arquivos MPG são arquivos de vídeo compactados, isto é, muito menores, em bytes, que seus equivalentes em AVI.

Os arquivos MPG são usados em diversos segmentos, como em DVDs (os DVDs de filme

mantêm seus conteúdos em arquivos de formato MPG, apesar de usarem outra extensão). Aliás, MPEG é o nome do grupo de profissionais (programadores) que cria e mantém os principais algoritmos (programas) para a compactação de sons e vídeos.

MPEG significa ***Motion Picture Experts Group*** (ou Grupo de Especialistas em Imagens em Movimento).

Há algumas variações para o MPG, como M4V, MP4, MPEG4 entre outros! Esses são todos arquivos de vídeo!

### **Extensão WAV (Som Real)**

Os arquivos com extensão WAV são arquivos de som. Esses arquivos não possuem nenhum nível de compactação e normalmente são muito grandes. Os arquivos WAV não oferecem nenhum tipo de habitat para a existência de vírus de computador.

Os sons gravados em formato WAV são reais, gravados por meio de microfones ou mesas de som, permitindo o armazenamento de todos os detalhes do som natural.

### **Extensão MP3 (Som Real Compactado)**

Como você já deve saber, os arquivos MP3 são usados para armazenar som, como os arquivos WAV. Os arquivos com essa extensão são compactados, ou seja, gastam menos bytes que os seus equivalentes em WAV.

Na verdade, os arquivos MP3 são arquivos WAV que passaram pelo processo de compactação. Esse processo é chamado de ***MPEG Layer 3***. Então, só para você lembrar: MP3 significa MPEG-3, ou MPEG Layer 3.

Como virou uma verdadeira “febre”, muitos equipamentos eletrônicos pessoais (DVD players, som automotivo, diskman, CD player, telefones celulares, gravadores de voz etc.) são capazes de ler os dados compactados no formato MP3, tocando o áudio deles. Em suma, um “equipamento MP3” nada mais é que um dispositivo qualquer com a capacidade de ler arquivos no formato MP3.

### **Extensão WMA (Som Real Compactado – Microsoft)**

Os arquivos com extensão WMA contêm som, como os arquivos WAV ou MP3. Esses arquivos são compactados, assim como os MP3, mas a principal diferença entre eles é o algoritmo (processo matemático) usado para a compactação.

WMA significa ***Windows Media Audio***. Isso indica que esses arquivos são criados e lidos pelo programa Windows Media Player, usando um algoritmo próprio da Microsoft, a produtora do Windows e de seus aplicativos.

Embora seja um formato exclusivo da Microsoft, o WMA já pode ser lido por diversos equipamentos e programas distintos. Uma das vantagens do WMA em relação ao MP3 é que o arquivo em WMA fica um pouco menor que o equivalente MP3.

### **Extensão WMV (Vídeo Compactado – Microsoft)**

São arquivos de vídeo (como os MPG) que podem ser lidos no programa Windows Media Player, da Microsoft. Esses arquivos são compactados através de um algoritmo próprio da Microsoft e normalmente não têm compatibilidade com nenhum outro programa de vídeo.

WMV significa, como você deve ter deduzido, Windows Media Video.

### **Extensão MOV (Vídeo do QuickTime)**

Os arquivos com extensão MOV são vídeos, criados e lidos pelo programa *Quick Time*, da Apple. Esses arquivos normalmente possuem uma qualidade de imagem e som muito boa e são usados amplamente pela indústria cinematográfica para disponibilizar, na Internet, trailers de filmes e outros produtos promocionais. Para abrir arquivos desse tipo, é necessário possuir o programa Quick Time, da empresa Apple.

### **Extensão MID (Som Sintetizado)**

Os arquivos com extensão MID (ou MIDI) são arquivos de som sintetizado, ou seja, criado pelo computador com notas musicais limitadas. Normalmente são sons na forma de apitos (tun, tuuuuun, TUN, tuuuuuun... notaram como é afinado?) que podem reproduzir, de forma bem artificial, músicas variadas. Esses são os sons usados nos toques (ringtones) dos telefones celulares (quando esses não usam MP3).

A principal diferença entre MID e WAV é que o segundo tipo guarda som real, como vozes, bateria, guitarra, gritos, rangidos de portas etc. O MID, por sua vez, reproduz apenas os apitos (tuuuuns) em diferentes frequências (representando, assim, notas musicais diferentes).

Os arquivos MID são muito menores que os arquivos WAV. (Isso se deve ao fato de eles não guardarem as mesmas coisas!)

#### **4.2.4. Atalhos**

Desde as primeiras versões (95 e NT), o sistema Windows apresentou um recurso bastante útil chamado atalho. Um atalho é somente um arquivo (com extensão LNK) que “aponta” para outro recurso qualquer acessível pelo Windows.

Um atalho pode apontar para outro arquivo qualquer, uma pasta, uma unidade de disco, um website, um arquivo localizado na Internet, uma impressora, outro computador na rede etc.

Portanto, quando se executa um atalho (duplo clique em seu ícone), na realidade, está se executando o recurso para o qual o atalho aponta (seja um site, um arquivo, uma pasta ou qualquer outro item).

Os atalhos são reconhecidos por uma característica visual peculiar: uma pequena setinha na extremidade inferior esquerda do ícone. Todo ícone que apresentar tal setinha é, na verdade, um atalho.





*Figura 4.5 – Atalho para o programa “Mozilla Firefox” entre alguns arquivos.*

Também são atalhos os ícones que ficam no Menu Iniciar (aqueles ícones que ficam no menu Programas).

Quando um atalho é apagado (excluído), apenas ele é afetado e não o seu alvo. Ou seja, se um atalho, localizado na área de trabalho do Windows (Desktop) e que aponta para o programa Word, for excluído, o programa Word nada sofrerá.

Se o arquivo para o qual o atalho aponta for excluído, o atalho também nada sofrerá; apenas, quando for clicado, não saberá para quem apontar, porque seu alvo terá sumido. Meio óbvio, não?

#### **4.3. Windows 7 – O mais atual**

No final de 2009, a Microsoft lançou mais uma versão do seu sistema operacional: o Windows 7. A principal característica desse sistema é o visual, que mudou radicalmente em comparação aos antecessores.

Atualmente, o Windows XP (lançado em 2001) ainda é o mais utilizado (e, com isso, mais cobrado em provas), mas, devido à aceitação, e ao tempo de lançamento, o Windows 7 tende a ser cada vez mais cobrado!



Figura 4.6 – A “cara” do Windows 7 – com três janelas abertas.

Há várias versões disponíveis do Windows 7:

- **Windows 7 Starter:** uma versão muito simplória e limitada. Não é vendida separadamente. Normalmente encontra-se incluída em hardwares (laptops e netbooks) quando os compramos.
- **Windows 7 Home Basic:** versão indicada para público doméstico. Normalmente é comprada por aqueles que compraram os hardwares citados acima e não aguentaram as limitações do Starter.
- **Windows 7 Home Premium:** para o público caseiro, é a melhor opção, pois traz mais recursos (programas) que a Home Basic.
- **Windows 7 Professional:** ideal para ser instalada em computadores de usuários em

empresas. Traz ferramentas (programas) especializados em trabalho em rede, por exemplo.

- **Windows 7 Ultimate:** a mais completa (e cara) versão do Windows 7. Traz todos os recursos que todas as versões anteriores trazem!

Mais informações acerca das comparações entre essas edições do Windows 7 podem ser conseguidas aqui (site oficial da Microsoft): <http://windows.microsoft.com/pt-BR/windows7/products/compare>

(Na boa, acho que não é necessário decorar o que tem lá, não!)

Lembre-se: nenhum dos produtos da família Windows é gratuito e tampouco código aberto, como o Linux. Bill Gates não seria um dos homens mais ricos do mundo se distribuisse programas de graça por aí, não é?

#### **4.3.1. Principais componentes do Windows 7**

##### **4.3.1.1. Desktop (Área de trabalho)**

É o nome dado à tela inicial do sistema operacional Windows. Todo usuário de computador que trabalha com o Windows conhece esta tela:



*Figura 4.7 – Desktop do Windows.*

#### 4.3.1.2. Barra de tarefas

É a barra horizontal que atravessa toda a base da área de trabalho. Essa barra apresenta o botão Iniciar, os botões dos Programas (fixos e abertos) e a Área de Notificação (onde está o relógio) e, na extremidade direita, um “botão” chamado Mostrar a Área de Trabalho (aquele retângulo que parece ser de vidro).



*Figura 4.8 – A barra de tarefas.*

#### **4.3.1.3. Botão Iniciar/Menu Iniciar**

É o botão que dá acesso a todos os recursos e programas no Windows.

Ao clicar no botão Iniciar, surge o **Menu Iniciar**, a partir de onde podemos iniciar qualquer programa, aplicativo, ou configuração que desejarmos no Windows. Na figura a seguir, o Windows 7 está apresentando seu Menu Iniciar.



Diablo III



Snagit 11 Editor



Google Chrome



Sothink DVD Ripper



Microsoft Excel 2010



Calculadora



Windows Live Mail



Microsoft Office Word 2007



vMix 2012 (x64)



Microsoft Office PowerPoint 2007



BrOffice Writer



Ferramenta de Captura

Todos os Programas

Joao Antonio

Documentos

Imagens

Músicas

Jogos

Computador

Painel de Controle

Dispositivos e Impressoras

Programas Padrão

Ajuda e Suporte

Desligar



*Figura 4.9 – O botão Iniciar dando acesso ao Menu Iniciar.*

Podemos destacar alguns dos principais componentes do Menu Iniciar do Windows a seguir:

### **Lista de Programas Fixos**

Os ícones que estão localizados no canto superior esquerdo do menu são atalhos que o usuário coloca por sua conta. Ali, o usuário poderá colocar atalhos para quaisquer programas que julgue interessantes (seus programas favoritos, por exemplo).

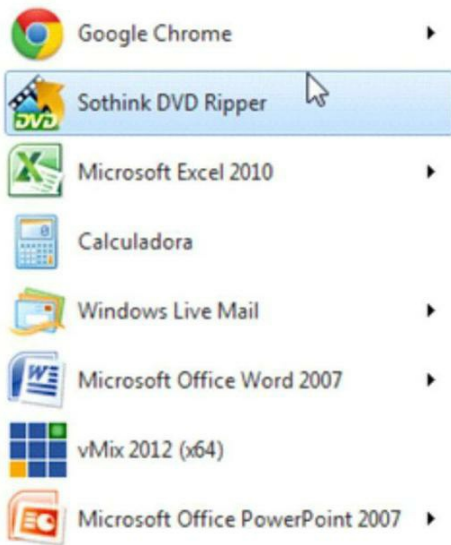
Note que na figura acima, dois programas estão fixos (ou seja, sempre estarão lá): o Diablo III (jogo que em que eu sou simplesmente viciado!) e o Snagit 11, programa que utilizo para capturar as telas que coloquei neste livro!



*Figura 4.10 – Programas fixos no Menu Iniciar.*

### **Lista de Programas Mais Usados**

Os ícones que ficam na parte inferior esquerda do menu são atalhos para os programas usados pelo usuário mais recentemente. (Essa listagem de atalhos muda constantemente de acordo com o que o usuário manipula no computador.)



*Figura 4.11 – Atalhos para os programas usados mais recentemente no Windows.*

É possível configurar quantos itens poderão ser apresentados nessa lista, bem como limpar a lista inteira para que ela seja preenchida conforme o usuário abra seus programas mais usados.

Perceba que alguns dos itens desta listagem (e dos itens fixos também) possuem uma pequena setinha apontando para a direita. Esta setinha é o indicativo de que o programa em questão foi usado para abrir vários arquivos recentemente.

Basta posicionar o ponteiro do mouse sobre o item (digamos o Microsoft Word 2007) que automaticamente a coluna à direita se transformará, mostrando a lista de arquivos mais recentemente manipulados pelo programa Word 2007 (conforme figura a seguir).



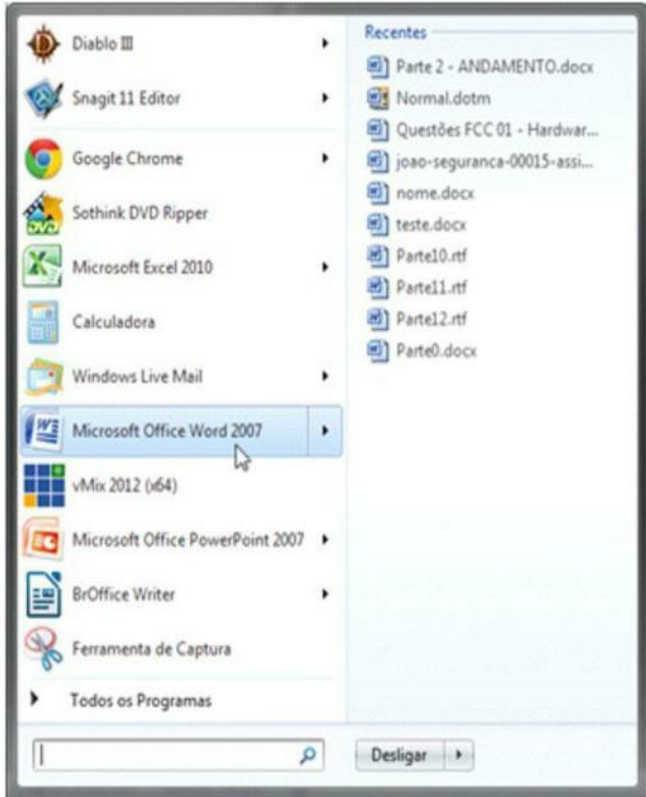


Figura 4.12 – Mouse posicionado sobre uma opção no Menu Iniciar.

## Principais Locais e Bibliotecas do Windows 7

Os ícones (atalhos) localizados na parte superior direita do Menu Iniciar apontam para os locais (pastas e bibliotecas) mais usados no Windows. Aqui vão explicações sucintas sobre os itens desta área do Menu Iniciar:

- **João Antonio (no seu micro, o nome é diferente):** é a pasta *pessoal* do usuário. Supostamente, é a pasta onde o usuário armazenará todos os seus arquivos de uso pessoal. Dentro desta pasta, no Windows 7, estão contidas as pastas “Documentos”, “Vídeos”, “Imagens” e “Músicas”, entre outras.

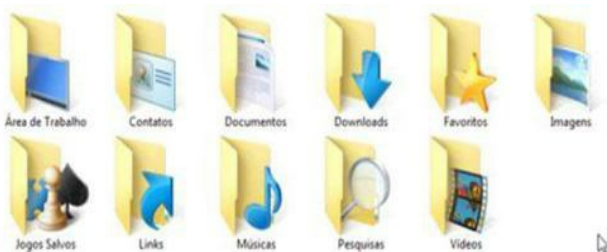
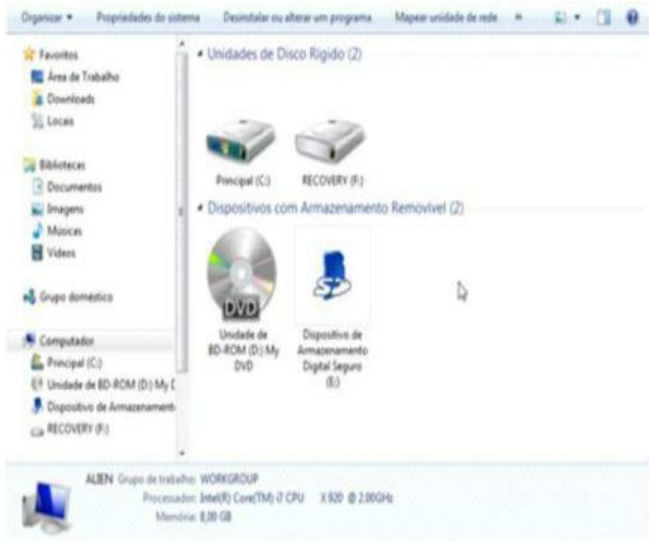


Figura 4.13 – Pastas contidas na Pasta Pessoal do usuário do Windows 7.

- **Documentos:** dá acesso à biblioteca Documentos (quase semelhante à pasta “Documentos” do usuário) – falaremos de bibliotecas posteriormente.
- **Imagens:** dá acesso à biblioteca Imagens (novamente: não é exatamente a mesma coisa que a pasta “Imagens” do usuário).
- **Músicas:** dá acesso à biblioteca Músicas.
- **Jogos:** dá acesso à pasta Jogos, que reúne todos os jogos instalados no Windows (pertencentes, ou não, ao sistema Windows).
- **Computador:** esse é o “local” principal do seu micro! Ele representa o seu micro! Dentro da janela “Computador”, estão listadas as unidades de disco presentes na sua máquina!



*Figura 4.14 – Janela “Computador”, mostrando as Unidades de Disco.*

### **Outros Comandos do Menu Iniciar**

Por fim, aqui vão os comandos restantes do Menu Iniciar:



*Figura 4.15 – Ícones restantes do Windows.*

- **Painel de Controle:** abre a janela do programa Painel de Controle, que reúne quase que a totalidade dos comandos de configuração do Windows. Vamos ver o Painel de Controle mais adiante!
- **Dispositivos e Impressoras:** dá acesso à janela que lista dispositivos de hardware (em sua maioria periféricos externos) ligados ao computador. Veja na figura a seguir:

#### Dispositivos (5)



ALIEN



Dispositivo de  
Entrada USB



Laptop\_Integrat  
ed\_Webcam\_2M



SD Memory Card



USB Receiver

#### Impressoras e Faxes (5)



Enviar para o  
OneNote 2010



Fax



HP Color Laserjet  
4700 PS



Microsoft XPS  
Document Writer



Snagit 11

Figura 4.16 – Janela dos “Dispositivos e Impressoras”.

• **Programas Padrão:** permite a configuração acerca de quais aplicativos (programas) serão usados para abrir quais arquivos (por sua extensão). Além de outras configurações, como visto a seguir:

## Escolha os programas que o Windows usa por padrão



### Definir os programas padrão

Tornar um programa o padrão para todos os tipos de arquivos e protocolos que ele pode abrir.



### Associar um tipo de arquivo ou protocolo a um programa

Fazer com que um tipo de arquivo ou protocolo (como .mp3 ou http://) sempre seja aberto em um determinado programa.



### Alterar configurações de Reprodução Automática

Reproduzir CDs ou outra mídia automaticamente



### Definir acesso a programas e padrões do computador

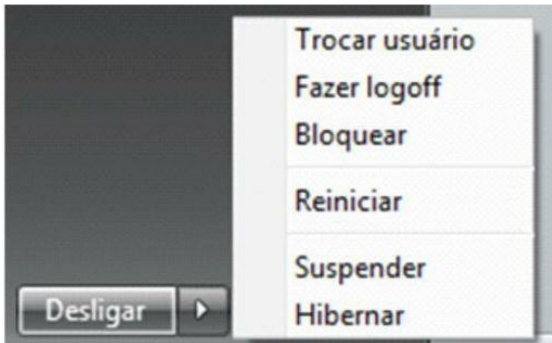
Controlar o acesso a certos programas e definir os padrões para este computador.

*Figura 4.17 – Janela dos “Programas Padrão”.*

- **Ajuda e Suporte:** permite acessar a janela de busca de ajuda (Help) do Windows 7. Esse sistema permite obter respostas às principais dúvidas do usuário tanto em textos presentes no próprio Windows como na base de dados da Microsoft na Internet.

## Comandos de Desligamento do Sistema

Na parte inferior do Menu Iniciar, encontra-se o Drop Down (caixa de listagem) dos comandos de Desligamento do Windows.



*Figura 4.18 – Opções de desligamento do sistema.*

Detalhes sobre as opções de desligamento (como Trocar usuário, Fazer logoff, entre outros) serão vistos no final deste capítulo.

### **Ferramenta para Pesquisar no Menu Iniciar**

Perceba que, na parte inferior do Menu Iniciar, bem à esquerda do botão de desligamento, há um campo de pesquisa, como mostrado na figura a seguir:



*Figura 4.19 – Campo de Pesquisa do Menu Iniciar.*

“Para que serve, João?”

Bem, caro leitor... Experimente-o! Você verá que ao digitar qualquer coisa neste campo, o Menu Iniciar se limitará a mostrar os itens que satisfazem aos critérios digitados por você neste campo!

### **Accionando o Menu Iniciar**

A maneira mais comum de abrir o Menu Iniciar é aplicar um clique no botão Iniciar, mas também é possível acionar a combinação de teclas CTRL+ESC para iniciar esse menu.

Ainda há, na maioria dos teclados, uma tecla específica para essa finalidade, com o formato do símbolo do Windows. Costuma-se chamá-la de *Tecla Windows*. Basta acioná-la uma única vez e o Menu Iniciar vai se abrir.



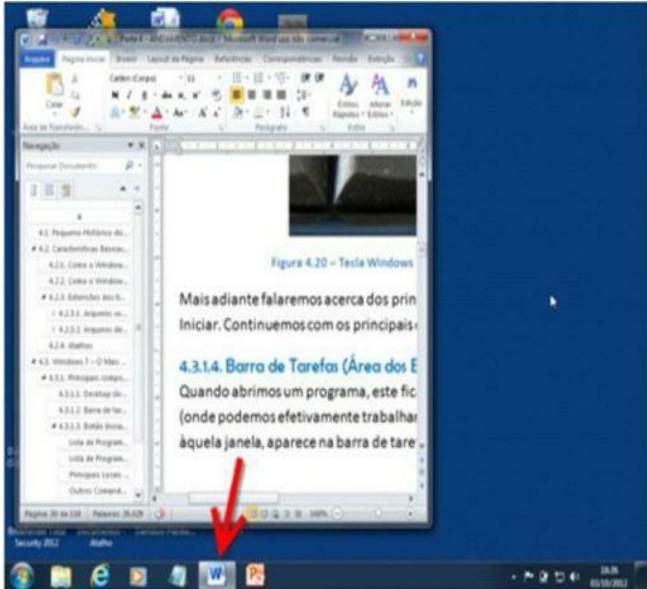
*Figura 4.20 – Tecla Windows (também chamada de Tecla Win).*

Mais adiante falaremos acerca dos principais componentes acessíveis pelo Menu Iniciar. Continuemos com os principais componentes do desktop.

#### **4.3.1.4. Barra de Tarefas (Área dos botões e programas)**

Quando abrimos um programa, este fica apresentado na forma de uma janela (onde podemos efetivamente trabalhar com ele) e um pequeno botão, referente àquela janela, aparece na barra de tarefas. Veja no exemplo a seguir:





*Figura 4.21 – Janela do Word e seu botão na Barra de Tarefas.*

Note, também, caro leitor, que há vários outros botões na Barra de Tarefas, mas eles se apresentam um pouco “diferentes” do botão do Word (dá para perceber que o botão do Word está “destacado”, como se tivesse uma borda ao seu redor e fosse feito de “vidro”, né?).

Pois é: os demais botões não possuem isso porque representam programas que não estão abertos (ou seja, não estão em funcionamento na memória RAM). A única forma, porém, de tais programas terem botões na Barra de Tarefas sem que estejam abertos é terem sido fixados lá!

Ou seja, a Barra de Tarefas apresenta botões para programas abertos naquele momento ou para programas que tiveram seus botões lá fixados.

Note a figura a seguir.



*Figura 4.22 – Botões da Barra de Tarefas.*

No exemplo acima, temos, em sequência: Windows Explorer, Internet Explorer, Windows Media Player e Bloco de Notas fixos na barra de tarefas. Além destes, os programas: Word, PowerPoint e Google Chrome abertos neste momento.

Note bem: sobre os programas abertos, não podemos afirmar se eles também estão fixos ou não! Ou seja, se um programa está aberto, e você percebe isso olhando para a barra de tarefas, você não poderá afirmar nada acerca de se ele está fixo ou não!

Perceba também que o último botão (o Google Chrome) está mais “claro”, mais “destacado” que os outros dois. Isso se dá porque o Google Chrome é a janela ativa (ou seja, ele é o programa que está à frente dos demais!).

Tais detalhes visuais só seriam cobrados em provas de bancas examinadoras que usam fotografias das telas do computador (até hoje, neste quesito, o Cespe/UnB é imbatível!).

Ahhh! Só lembrando... Quando abrimos alguns programas, pode ser que estes apresentem botões da seguinte forma (repare no botão do Word):



*Figura 4.23 – O Word com várias janelas abertas.*

Essa é a indicação visual de que há várias janelas do Word abertas simultaneamente no computador. Se você forçar a barra, verá que são três (pelas três bordas mostradas à direita do botão).

Se você clicar num botão assim (com várias janelas abertas), o Windows 7 lhe mostrará um painel contendo as miniaturas das várias janelas, permitindo que você escolha, no clique, qual delas trará para a frente!

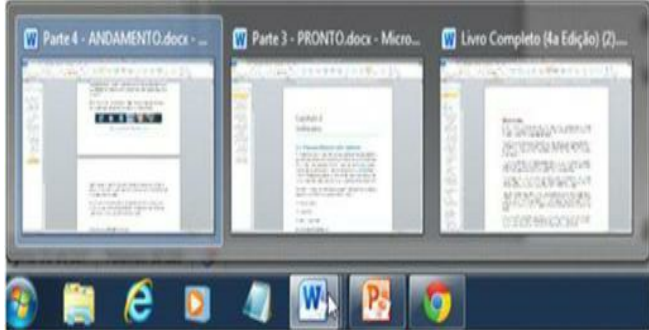


Figura 4.24 – Painel mostrando as três janelas do Word abertas.

Há muito mais operações que podem ser realizadas com os botões da Barra de Tarefas! Como são muito simples de fazer, mas demorariam muitas páginas para descrever (o que faria o livro ficar maior), sugiro que você busque esse material em vídeo, tanto no site da **Editora Campus/Elsevier** ([www.elsevier.com.br](http://www.elsevier.com.br)), na seção referente a este livro, quanto no **Eu Vou Passar** ([www.euvoupassar.com.br](http://www.euvoupassar.com.br)), no curso completo de Windows 7).

#### 4.3.1.5. Área de notificação (System Tray)

É a área à direita da Barra de Tarefas que apresenta o relógio do computador e outros ícones de programas em execução, como antivírus e outros programas residentes na memória.



Figura 4.25 – Área de notificação do Windows.

A grande maioria dos ícones apresentados na bandeja do sistema (ou system tray – o outro nome da área de notificação) representa programas executados em segundo plano, ou seja, sem a interferência do usuário. São programas que estão em funcionamento, portanto consomem memória RAM.

Novos ícones não podem ser colocados aqui pelo usuário, mas quando certos programas são instalados, eles mesmos tratam de se colocar nessa área.

É possível reconhecer alguns ícones básicos, pertencentes ao próprio Windows, como o:

- ***Alertas do Sistema (Central de Ações) – o ícone da bandeirinha:*** apresenta problemas do sistema e sugere soluções, bem como permite abrir a central de ações do Windows, que centraliza a tomada de providências para a manutenção do bom funcionamento do Windows.
- ***Energia – o ícone da “bateria” com a tomada:*** dá acesso às opções de energia do computador. Esse ícone normalmente só aparece quando se trata de um micro portátil (laptop, netbook, ultrabooks etc.).
- ***Rede – o ícone do monitor com o cabo de rede:*** indica que há acesso a uma rede de computadores com cabo (desenho do cabo, né?). Esse ícone dá acesso às informações pertinentes à conexão com a rede, bem como permite configurá-las!
- ***Controle de Volume – o ícone do alto-falante:*** dá acesso aos controles de volume de som (tanto do som que se ouve nas caixas de som, como dos sons que entram no computador por meio dos microfones).

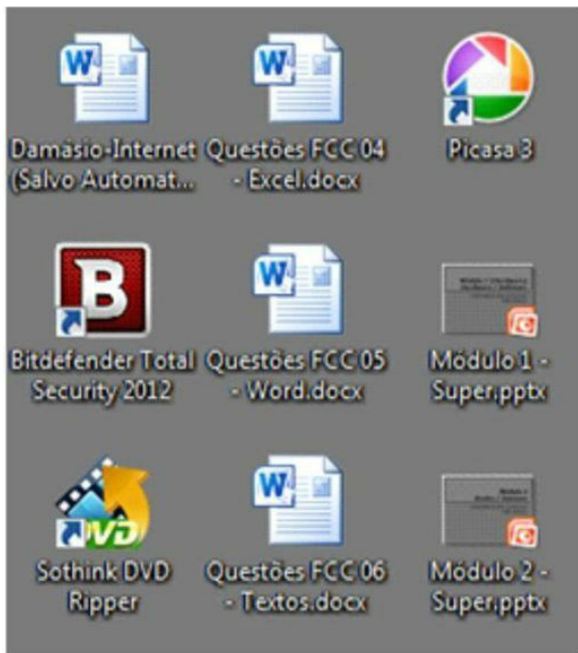
Além destes ícones, é possível verificar a existência de uma pequena setinha apontando para cima na extremidade esquerda desta área. Esta setinha, quando clicada, permite o acesso aos demais programas residentes (programas que ficam abertos sempre) na memória.



Figura 4.26 – Demais ícones da área de notificação.

#### 4.3.1.6. Ícones

São todos os pequenos símbolos gráficos que representam objetos utilizáveis no Windows. São os ícones que, quando abertos, iniciam programas, jogos, documentos etc. Na área de trabalho (desktop) do Windows, há vários ícones já colocados pelo próprio sistema e outros que o usuário pode colocar para facilitar sua vida (os atalhos).



Um ícone pode ser aberto (executado) de várias maneiras:

1. Aplicando um duplo clique nele.
2. Clicando uma vez nele (para selecioná-lo) e pressionando a tecla ENTER.
3. Clicando no mesmo com o botão direito do mouse e acionando o comando Abrir.

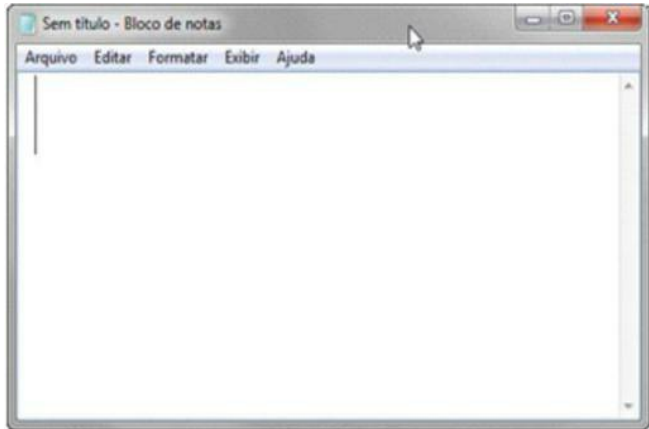
#### **4.3.1.7. Janelas**

Quando um ícone é aberto, ele se transforma em uma janela. Basicamente, todos os programas em execução (em funcionamento) são apresentados como janelas.



Figura 4.28 – Janela do “Computador” (programa Windows Explorer).

Algumas janelas apresentam itens diferentes, com formatos diferentes. Vamos nos ater, primeiramente, aos componentes mais “tradicionais” das janelas, como as que vamos ver agora (no exemplo, a janela do programa Bloco de Notas).



*Figura 4.29 – Janela do Bloco de Notas.*

#### **4.3.2. Componentes de uma janela**

A seguir, veremos os componentes que formam uma janela.

##### **4.3.2.1. Barra de título**

É a barra superior da janela (ela apresenta o nome do programa e/ou documento que está aberto). Apresenta em sua extremidade esquerda o Ícone de Controle e à direita os botões Minimizar, Maximizar e Fechar.



*Figura 4.30 – Barra de título da janela.*

Só para lembrar, caro leitor: nem todas as janelas possuem barra de título exatamente como esta. (dá para perceber isso na figura 4.28, mais acima).



#### 4.3.2.2. Barra de menus

É o conjunto de menus (listas de opções) dispostos horizontalmente, abaixo da barra de título. Cada item desta barra pode ser aberto com um único clique.

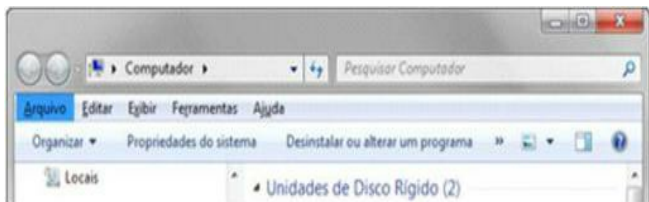
**Arquivo    Editar    Formatar    Exibir    Ajuda**

*Figura 4.31 – Barra de menus.*

Algumas janelas parecem não possuir barra de menus (alguns poucos programas realmente não a tem). Mas a maioria simplesmente “aparenta” não possuir, ou seja, eles têm a barra de menus, mas normalmente não mostram isso!

Quando o usuário pressiona a tecla ALT, cada item do menu apresenta uma de suas letras sublinhada. Basta acionar a letra sublinhada no menu desejado (ainda com a tecla ALT pressionada) e esse menu se abrirá. Exemplo: o menu Arquivo fica com a letra A sublinhada; portanto, a combinação de teclas ALT+A serve para abri-lo.

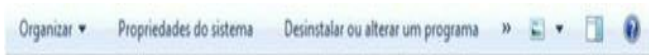
Nas janelas onde a barra de menus é escondida, pressionar a tecla ALT faz com que ela apareça, conforme se pode ver na figura a seguir:



*Figura 4.32 – Janela do Computador mostrando a barra de menus.*

#### 4.3.2.3. Barra de ferramentas

É a barra horizontal que apresenta alguns botões de comandos que acionam comandos existentes nos menus. Os comandos desta barra normalmente existem nos menus, mas é mais rápido executá-los por esse meio.



*Figura 4.33 – Barra de ferramentas.*

É bom lembrar que nem todas as janelas e/ou programas apresentam a barra de ferramentas. Pode-se ver claramente que o Bloco de Notas (Figura 4.29) não a possui!

#### 4.3.2.4. Barra de endereço

Apresenta o endereço do local cujo conteúdo está sendo visualizado na janela. No nosso caso, estamos visualizando o conteúdo do item **Computador**. Nesse local, podemos digitar um endereço de uma pasta do seu computador, uma unidade de disco, outro computador da rede ou até mesmo um site da Internet.

No Windows Explorer do Windows 7, a barra de endereços é “mesclada” com a barra de título (é tudo uma coisa só!) – e a barra de endereço ainda possui um “segredinho” novo (que os Windows anteriores não possuíam).

Vamos abordar isso no tópico que fala do Windows Explorer, mais adiante!



*Figura 4.34 – Barra de endereço + Campo de pesquisa*

Algumas janelas trazem, ao lado da barra de endereço, um campo para fazer pesquisas no computador, conforme mostrado acima! Vamos conhecer todos esses detalhes no Windows Explorer, mais para a frente!

#### 4.3.2.5. Barra de status

Apresenta algumas informações sobre o conteúdo da janela em questão. Atenção, pois a barra de status é um dos mais importantes componentes das janelas, já que pode trazer uma série de informações interessantes para resolver questões de provas (especialmente, aquelas que utilizam fotos das janelas).



### 4.3.3. Principais operações com janelas

#### 4.3.3.1. Movendo uma janela

Para mover uma janela (alterar sua posição na tela), basta clicar na barra de título da janela e arrastá-la até a posição desejada. Observe que o arrasto tem de ser feito pela barra de título da janela.

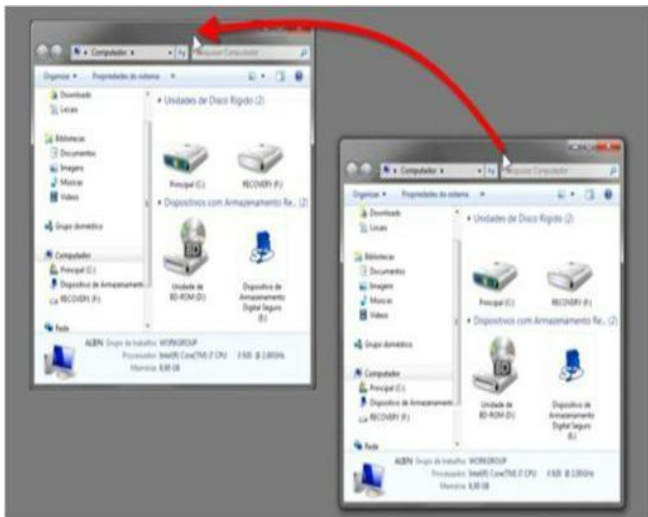
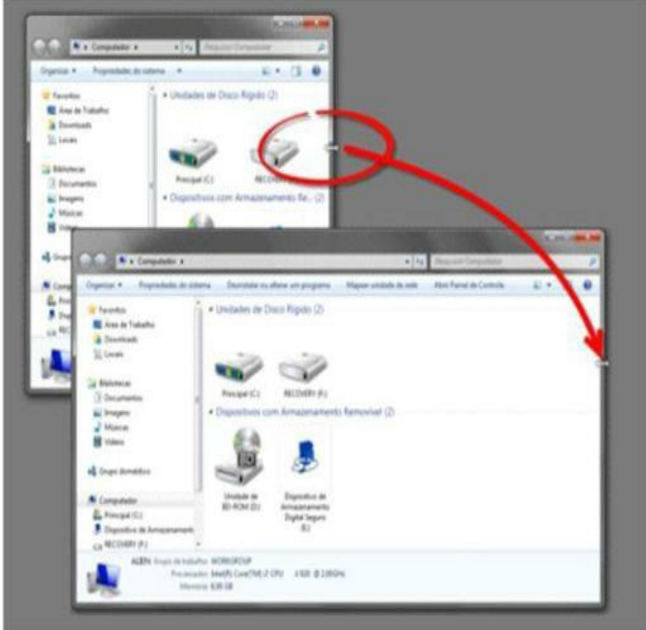


Figura 4.36 – Para mover uma janela, deve-se arrastá-la pela barra de título.

#### 4.3.3.2. Redimensionando uma janela

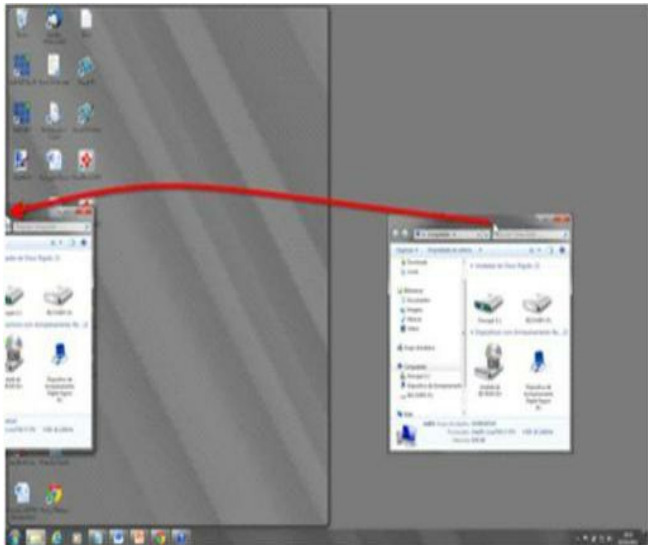
Redimensionar uma janela significa alterar seu tamanho (largura ou altura). Para fazer isso, basta clicar em uma das bordas da janela (o ponteiro do mouse se transformará em uma seta dupla) e arrastar até a forma desejada para a janela.



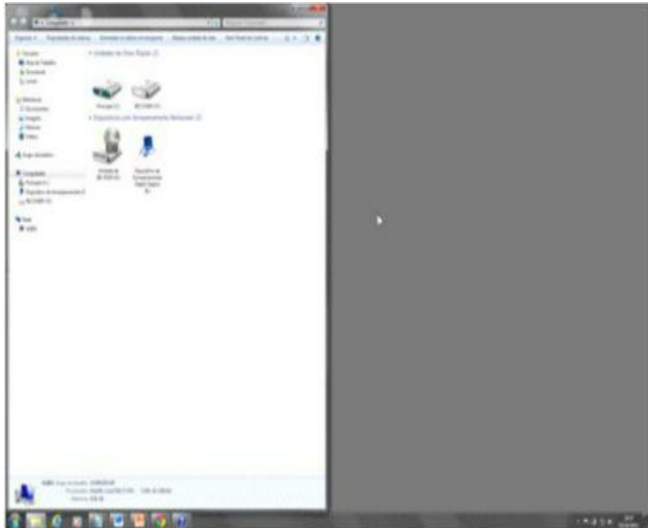
*Figura 4.37 – Redimensionando uma janela.*

O usuário poderá usar uma das bordas laterais (esquerda ou direita) para fazer o dimensionamento horizontal (como mostra a figura anterior) ou pode usar as bordas superior e inferior para um dimensionamento vertical. Posicionar o ponteiro do mouse em um dos quatro cantos (diagonais) da janela permite o dimensionamento livre (horizontal e/ou vertical simultaneamente).

Outra forma de dimensionar uma janela para que ela fique exatamente com o tamanho equivalente à metade do tamanho total da tela (sim, exatamente MEIA TELA) é arrastá-la (pela barra de título) até uma das extremidades laterais da tela (direita ou esquerda).



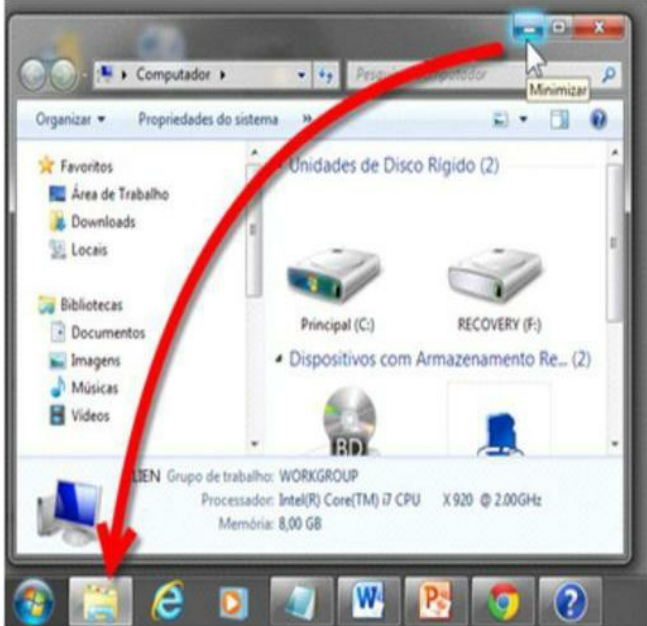
*Figura 4.38 – Arrastando a janela para a lateral esquerda.*



*Figura 4.39 – Janela já dimensionada (metade da tela).*

#### **4.3.3.3. Minimizando uma janela**

Minimizar uma janela significa fazê-la recolher-se ao seu botão presente na barra de tarefas do Windows. Para minimizar uma janela, basta clicar no botão Minimizar, em sua barra de título.



*Figura 4.40 – Clicar no botão Minimizar de uma janela a faz recolher-se ao botão.*

Para fazer a janela minimizada voltar a aparecer (chamamos isso de restaurar), basta um clique simples no seu botão correspondente na barra de tarefas. Também é possível minimizar uma janela clicando diretamente nesse botão (na barra de tarefas), se for a janela ativa (aquela que está na frente das demais janelas).

Há também formas de minimizar todas as janelas abertas de uma única vez. Para fazer isso, basta acionar o botão Mostrar Área de Trabalho, que fica na extremidade direita da barra de tarefas (ao lado do Relógio).



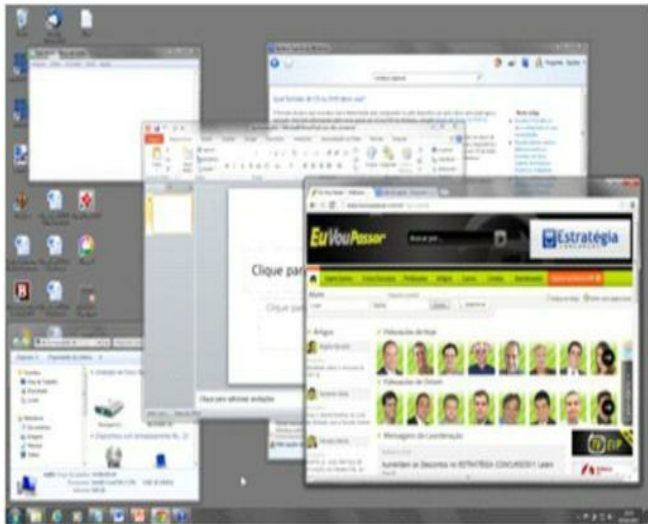
*Figura 4.41 – Botão Mostrar a Área de Trabalho.*

Se você posicionar o mouse sobre este botão (sem clicar), todas as janelas ficam momentaneamente transparentes (como se fossem de vidro) – neste momento, tirar o mouse fará as janelas voltarem a ser opacas. E se você clicar no botão, as janelas realmente minimizam, mostrando a área de trabalho permanentemente.

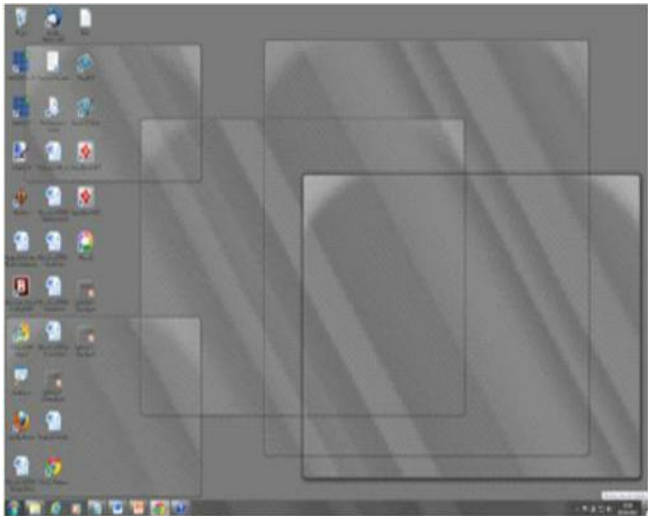
Um novo clique neste botão (quando todas as janelas estiverem minimizadas) fará com que todas as janelas sejam restauradas a seus estados originais.

Veja na sequência de figuras a seguir:





*Figura 4.42 – Todas as janelas abertas.*



*Figura 4.43 – Janelas translúcidas, enquanto o mouse está no botão Mostrar Área de Trabalho.*

Para minimizar todas as janelas, é possível, também, acionar a combinação de teclas **WINDOWS + M** (isso minimiza todas, mas não consegue restaurá-las) ou **WINDOWS + D** (isso minimiza todas e, se acionado novamente, restaura todas a seus estados originais).

#### **4.3.3.4. Maximizando uma janela**

Maximizar uma janela é fazê-la redimensionar-se para tomar todo o espaço possível da tela. Para fazer isso, clique no botão Maximizar.



*Figura 4.44 – Botão Maximizar.*

Quando uma janela é maximizada, o botão maximizar é substituído pelo botão Restaurar Tamanho, que faz a janela retornar ao tamanho que tinha antes da maximização.



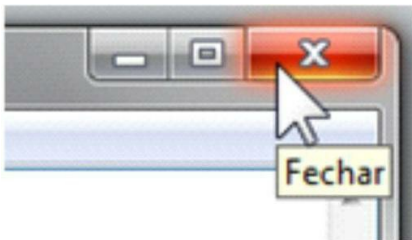
*Figura 4.45 – Botão Restaurar Tamanho.*

Maximizar (e restaurar abaixo) uma janela pode ser feito aplicando-se um clique duplo na barra de títulos da janela.

Também é possível maximizar uma janela (novidade no Windows 7) arrastando-a, por meio da barra de título, para a extremidade superior da tela (o movimento é semelhante ao redimensionar para ½ tela, que já vimos... Apenas com a diferença que a gente deve arrastar para CIMA!).

#### **4.3.3.5. Fechando uma janela**

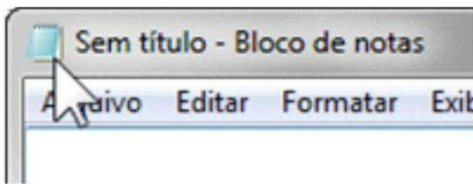
Para fechar uma janela, clique no botão Fechar (X) no canto superior direito da janela. Essa ação fará as informações da janela serem retiradas da memória RAM do computador e, com isso, o programa associado àquela janela será fechado.



*Figura 4.46 – Botão Fechar.*

Outra maneira de fechar uma janela é acionando a combinação ALT+F4 no teclado do computador. Essa ação fechará apenas a janela que estiver com o foco (janela ativa).

Também é possível solicitar o fechamento da janela ao acionar um DUPLO CLIQUE no Ícone de Controle da janela (o pequeno ícone que fica localizado na extremidade esquerda da barra de título).



*Figura 4.47 – Ícone de Controle; duplo clique fechará a janela.*

Através desse ícone também é possível acionar os outros comandos vistos até aqui, basta aplicar um CLIQUE SIMPLES no mesmo e o menu se abrirá.



*Figura 4.48 – Um único clique abrirá este menu.*

Outra forma de abrir o menu de controle da janela é acionando a combinação de teclas ALT + BARRA DE ESPAÇO.

#### **4.3.3.6. Trabalhando com várias janelas abertas**

Podemos abrir diversas janelas ao mesmo tempo no Windows, embora só seja possível manipular uma delas por vez. Para alternar entre janelas abertas, passando o foco de uma para outra, basta acionar ALT + TAB.

Bom, a forma certa de usar é segurar a tecla ALT e, mantendo-a pressionada, acionar TAB tantas vezes quantas forem necessárias até o foco estar na janela que se deseja trazer para a frente. Você poderá ver quem está “com o foco” por meio da pequena janela que aparece enquanto a tecla ALT está pressionada.

Quando você segura ALT e pressiona o TAB pela primeira vez, o painel abaixo é mostrado. A cada TAB que você pressionar posteriormente, a próxima miniatura de janela é selecionada! Ao soltar o ALT, é justamente esta janela que virá para a frente!



*Figura 4.49 – Painel da alternância entre janelas (ALT + TAB).*

Outra forma de alternar diretamente entre as janelas abertas (sem que se abra o painel mostrado acima) é através das teclas ALT + ESC (alternância direta)!

A terceira, e mais “enfeitada” forma de alternância é chamada de **Flip 3D**. É conseguida segurando a tecla WINDOWS e acionando TAB várias vezes (ou seja, WINDOWS + TAB).



*Figura 4.50 – Flip 3D – alternância com “estilo”.*

#### 4.4. Principais programas do Windows

Caro leitor, agora que você foi apresentado aos conceitos básicos do sistema operacional Windows (mais precisamente, na versão 7), é hora de conhecer os programas que o acompanham e suas principais funções.

##### 4.4.1. Windows Explorer

O Windows Explorer é o programa **gerenciador de arquivos** do sistema operacional Windows. É através do Windows Explorer que podemos manipular os dados gravados em nossas unidades,

copiando, excluindo, movendo e renomeando os arquivos e pastas das nossas unidades de armazenamento.

Sem dúvida alguma, o Windows Explorer é a mais importante ferramenta pertencente ao Windows cobrada em provas! Se há um único assunto a ser estudado sobre Windows, este assunto é o Windows Explorer.

#### 4.4.1.1. Conhecendo a interface do Explorer

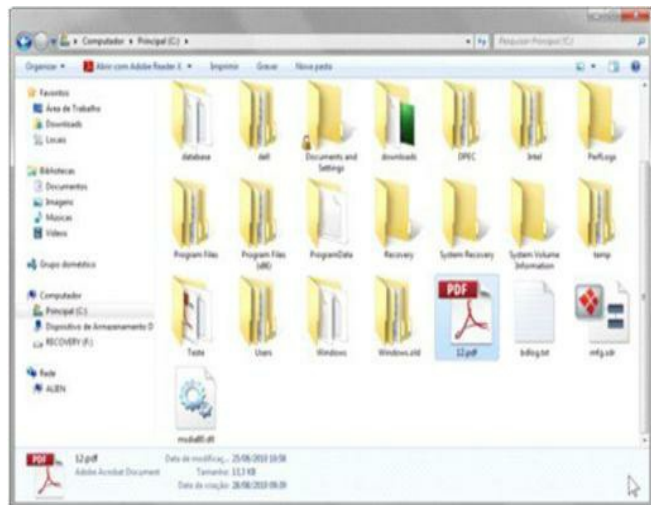
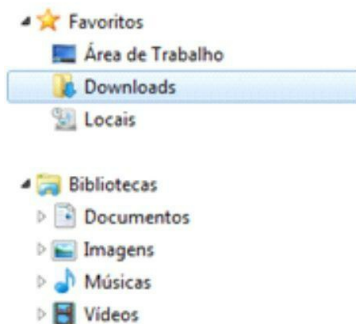


Figura 4.51 – O Windows Explorer.

O Windows Explorer apresenta sua interface dividida em duas partes: o painel da navegação (ou área das “pastas”), localizado à esquerda da janela, e o painel do conteúdo (a área grande à direita).

O Painel de Navegação (também chamado de “área da árvore”) é o painel que mostra a estrutura completa do computador, hierarquicamente, pasta por pasta, unidade por unidade, como um grande organograma. Na área das pastas não são mostrados arquivos.

O Pannel de Navegação também mostra Bibliotecas (falaremos sobre elas mais adiante), lista de locais favoritos (no topo) e acesso aos computadores da rede, seja diretamente (mostrado na parte de baixo), seja pelo Grupo Doméstico (vamos citar este recurso interessante também!).



*Figura 4.52 – Detalhe do Pannel de Navegação.*

A área do conteúdo apresenta o que há na pasta selecionada da árvore. Na área de conteúdo pode aparecer todo tipo de objeto (arquivos, pastas, unidades). Ou seja, quando se clica em algum item no Pannel de Navegação, automaticamente seu conteúdo é mostrado no Pannel de Conteúdo.

No Windows Explorer, sempre deve haver um local explorado, ou seja, o programa sempre estará visualizando o conteúdo de algum diretório (pasta ou unidade). Para escolher o diretório cujo conteúdo será mostrado, basta clicar nele na árvore.

Na figura seguinte, é possível ver o usuário escolhendo uma pasta para visualizar seu conteúdo. Lembre-se: apesar de o clique ter sido dado no pannel de navegação (à esquerda), o conteúdo será mostrado no pannel à direita (área do conteúdo).





*Figura 4.53 – Um clique na Biblioteca Imagens permite ver seu conteúdo.*

Para saber qual o local (pasta, biblioteca, unidade etc.) que está sendo explorado no momento (que é interessante para as provas de concurso que apresentam fotografias, como as do Cespe/UnB), basta ler na barra de endereços do programa.

No caso da figura a seguir, estamos explorando uma pasta chamada **Support**, localizada dentro da pasta **ATI**, que, por sua vez, se encontra dentro da Unidade (C:). Simples, não?

Tudo isso pode ser lido, simplesmente na barra de endereços, localizada na barra de título da janela do Windows Explorer! Não se esqueça disso! Essa barra (endereços) é muito importante!

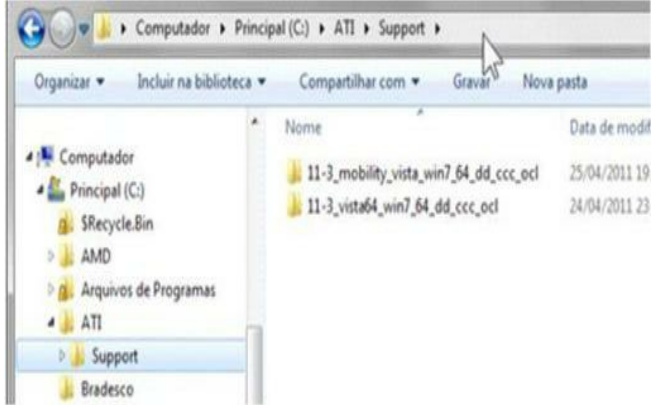


Figura 4.54 – Identificando o local que está sendo explorado.

### Entendendo a Barra de Endereços do Windows 7

No Windows Explorer do Windows 7, a barra de endereços trouxe algumas diferenças em relação a suas versões anteriores.

Em primeiro lugar, as “setinhas”. Perceba que a cada “novo nível”, há uma setinha entre ele e o anterior. Essa setinha não é só a indicação de que há níveis (do tipo “um dentro do outro”) entre aqueles locais. Mas vamos começar vendo isso dessa forma:



Figura 4.55 – A barra de endereços.

Sabe o que ela significa?

1. Estamos, neste momento, explorando a pasta **Quarto 2**. Dica: sempre estamos explorando o último nome mostrado na barra! Ou seja, o último nome mostrado na barra de endereços é, sem dúvidas, o nome da pasta (do local) que estamos explorando naquele momento.
2. **Quarto 2** está dentro de **Casa**. **Casa**, por sua vez, dentro da **Unidade de disco F:**, que, como toda unidade, está dentro do item **Computador**.

3. Quarto 2 tem subpastas (ou seja, existem pastas dentro da pasta Quarto 2) – Ahhh! Por essa você não esperava, né? Olha o detalhe: se o último nome (que, já sabemos, indica a pasta em que estamos) estiver seguido de uma setinha (e tá lá!), é sinal de que a pasta em questão (Quarto 2), tem subpastas (pastas dentro dela!). Não levou fé? Olha a foto a seguir!

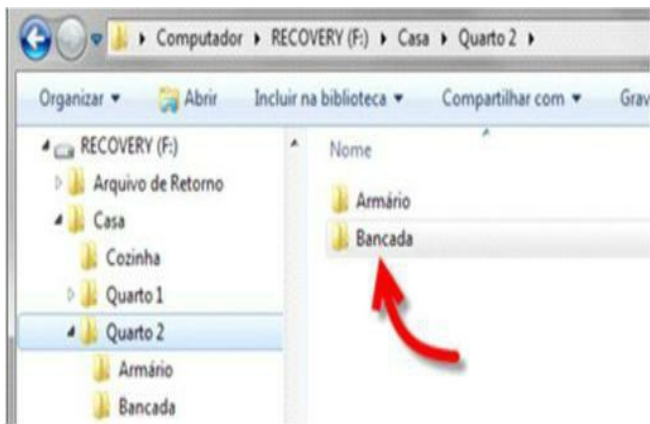


Figura 4.56 – Armário e Bancada são subpastas de Quarto 2.

Isso pode ser visto tanto no Painel de Navegação (com as pastas *Armário* e *Bancada* sendo vistas abaixo – subordinadas – da pasta Quarto 2) quanto no próprio Painel do Conteúdo, que mostra as duas, provando que são conteúdo (estão dentro) da pasta Quarto 2.

As setinhas são botões que permitem acessar as subpastas de qualquer item presente na barra de endereços! Por isso é que eu digo: “se é seguido por uma setinha, tem subpastas!”.

Para exemplificar o uso das setinhas para navegar entre subpastas: estamos na pasta Quarto 2, mas se quisermos ir para a pasta Cozinha (que é subpasta de Casa, assim como Quarto 2), basta clicar na setinha após Casa e escolher Cozinha na lista! Sim! Sim! Clicar na setinha! Saca só:

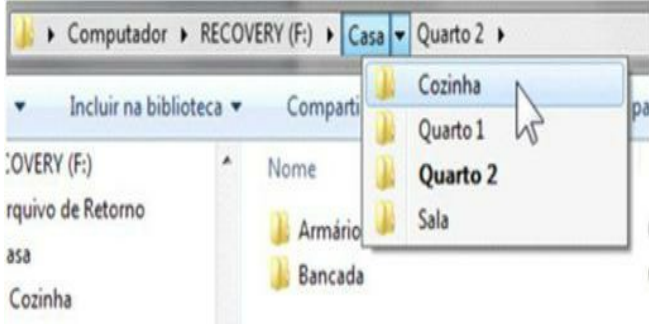


Figura 4.57 – Usando a “setinha” para acessar outras subpastas.

Note, também, que cada nome na barra de endereço é um botão, em si, que, se clicado, leva para aquela pasta em questão (ou seja, clicar em **Casa**, leva você à pasta **Casa**!). As setinhas, por sua vez, também são botões, mas que levam às subpastas daquela pasta anterior a elas!

Se você clicar no ícone que aparece à esquerda da barra de endereços, a barra, em si, passará a apresentar o endereço em questão de um “jeito antigo”, como era normal se apresentar no DOS e nos Windows (até o XP).

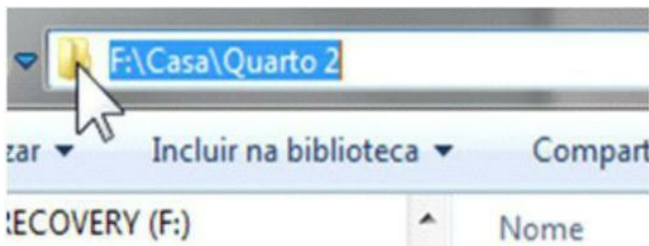


Figura 4.57a – Endereço em seu “formato tradicional”.

O formato que sempre foi usado no Windows para endereçar pastas é sempre este: o endereço

sempre *inicia com a unidade de disco* em questão e segue “*entrando*” em *cada pasta*, separando, sempre, os diretórios (pastas) dos subdiretórios (suas subpastas) por meio do sinal de \ (barra invertida, ou contrabarra).

Então, só como mais um exemplo, uma pasta chamada **Porta-luvas**, dentro de uma pasta chamada **Carro**, que está, por sua vez, dentro uma pasta chamada **Garagem**, que fica dentro da unidade C:, seria referenciada por meio do endereço **C:\Garagem\Carro\Porta-luvas**.

### Os Endereços na Árvore

Note, abaixo, a “árvore” que indica o endereço F:\Casa\Quarto 2\Bancada. Sabemos que Bancada é subdiretório (subpasta) de Quarto 2. Esta, por sua vez, é subpasta de Casa. Casa, por fim, está dentro da unidade F:.

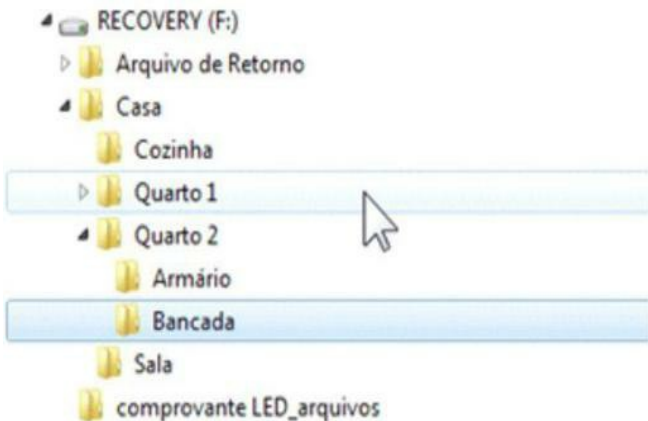


Figura 4.58 – O que significa F:\Casa\Quarto 2\Bancada.

O termo subdiretório (ou subpasta) é usado para definir uma relação entre um diretório e o seu nível imediatamente superior. No caso da figura anterior, Bancada e Armário são subdiretórios (subpastas) da pasta Quarto 2.

Note ainda que algumas pastas apresentam, à sua esquerda, um triângulo (que pode ser branco

ou preto) e outras simplesmente não apresentam tais sinais. As pastas que possuem triângulo possuem subpastas, já as pastas que não possuem triângulo não possuem subpastas (mas não podemos afirmar que estarão vazias, porque podem conter arquivos!!! Fique ligado!).



*Figura 4.59 – Pastas que contêm (com sinais) e não contêm (sem sinais) subpastas.*

Um clique no triângulo branco fará a pasta em questão ser expandida na própria árvore, mostrando suas subpastas (e o triângulo se transforma no preto). Quando se clica no triângulo preto, este fará a pasta em questão ser contraída, escondendo novamente suas subpastas na árvore (e ele volta a ser o triângulo branco).

Note: “expandir” e “contrair” são os verbos utilizados para descrever o ato de “mostrar” ou “esconder” as subpastas de uma determinada pasta em questão. Mostrando (expandir) suas ramificações ou escondendo-as (contraíndo).

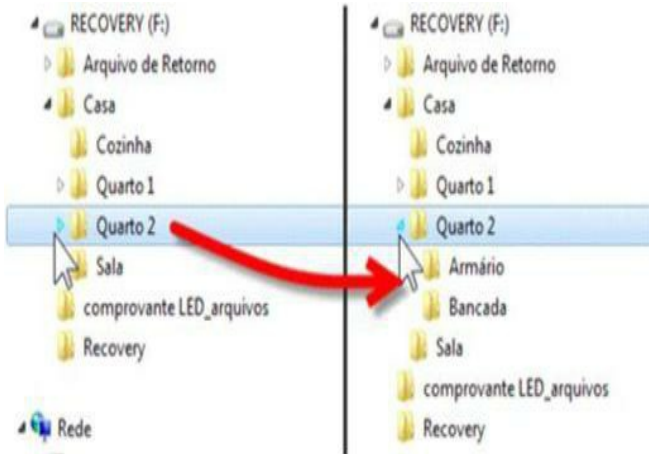


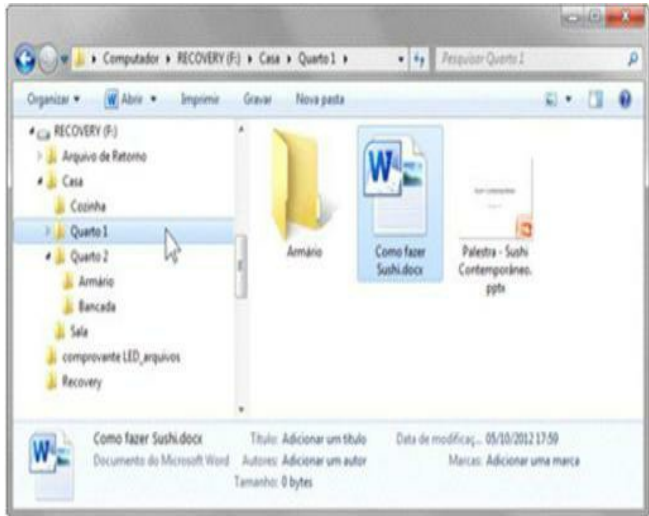
Figura 4.60 – Expandir versus Contrair...

No exemplo da figura anterior, as pastas Arquivo de Retorno, Casa, Quarto 1 e Quarto 2 possuem subpastas (Quarto 2 foi expandida no exemplo, mostrando suas subpastas). As demais pastas do exemplo não possuem subpastas.

Expandir e Contrair são ações que são realizadas e acontecem apenas no Pannel de Navegação (ou seja, apenas na parte esquerda do Windows Explorer). Quando usamos a expressão “*Abrir*” ou “*Explorar*”, consiste em dizer que a pasta está sendo visualizada, ou seja, que o seu conteúdo está sendo visto (isso, claro, acontece no Pannel do Conteúdo).

### Outros Detalhes da Interface (Área do Conteúdo)

Veja, a seguir, o conteúdo da pasta Quarto 1. Ou seja, neste momento, a pasta Quarto 1 está sendo explorada!



*Figura 4.61 – Pasta Quarto 1 contém uma pasta e dois arquivos.*

Os ícones do Painel de Conteúdo podem ser apresentados de várias formas, basta clicar no botão Modo de Exibição, localizado à direita na barra de ferramentas. Veja todas as opções:



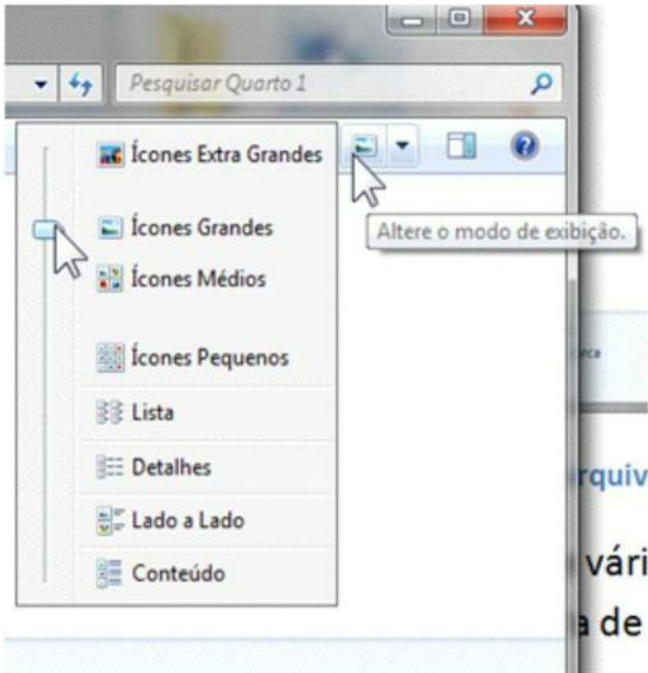


Figura 4.62 – Botão para alteração do modo de exibir os ícones.

Um dos formatos mais interessantes é o **Detalhes**, que mostra os objetos em lista vertical, acompanhados de várias informações interessantes sobre eles (das quais, claro, se pode extrair inúmeras questões de prova!).

| Nome                                                                                                                | Data de modificaç... | Tipo                | Tamanho |
|---------------------------------------------------------------------------------------------------------------------|----------------------|---------------------|---------|
|  Armário                             | 05/10/2012 17:16     | Pasta de arquivos   |         |
|  Como fazer Sushi.docx               | 05/10/2012 17:59     | Documento do Mi...  | 0 KB    |
|  Palestra - Sushi Contemporâneo.pptx | 05/10/2012 18:00     | Apresentação do ... | 37 KB   |

*Figura 4.63 – Exibição em modo Detalhes.*

Cada modo de exibição, porém, tem seu próprio “charme” e sua própria característica (que, diga-se de passagem, é assunto de questões de prova, também!). Visite-os, teste-os! Vai ser enriquecedor! Se quiser mais dicas, claro (macetes), visite o site da Editora Elsevier, lá você encontrará um vídeo explicativo sobre este conteúdo, especialmente preparado para você, leitor!

Aproveitando que estamos naquela “parte” da janela, eis que se apresenta o botão “Mostrar Painel de Visualização”, que permite ligar/desligar o terceiro painel do Windows Explorer: o painel que permite visualizar rapidamente o conteúdo de um arquivo selecionado na área de conteúdo.

Sem precisar abrir o arquivo para ver seu conteúdo (parte dele), é só clicar no arquivo e o painel de visualização (se estiver aberto, claro!) mostrará a primeira página, o primeiro slide, os primeiros trechos do arquivo em questão!



Figura 4.64 – Painel de visualização em ação.

O último item que é necessário conhecer acerca da “cara” do Windows Explorer é a barra de status, ou, hoje, Painel de Informações. É a área que fica na parte de baixo da janela, mostrando informações sobre o arquivo que estiver selecionado! Prepare-se para prestar bem atenção a ela, tá? Daqui, saem muitas questões de prova! O\_o



Figura 4.65 – Painel de Informações do Windows Explorer.

Além das informações naturais normais do arquivo (como tamanho, tipo, data de modificação), o Painel de Informações traz uma coisinha nova: as **Marcas**. Marcas são “palavras-chave” relacionadas ao arquivo, como “termos” que indicam a que o arquivo está associado.

Adicionar marcas ao arquivo poderá facilitar o agrupamento e a pesquisa de arquivos no seu computador. Basta clicar onde está “**Adicionar uma marca**”, e o campo se abrirá! Você poderá colocar diversas marcas, separadas por ponto e vírgula. Veja a seguir as marcas sendo

adicionadas ao arquivo:



Figura 4.66 – Adicionando uma marca.

Note que os campos **Título** e **Autores** também poderão ser alterados: experimente clicar neles para testar!

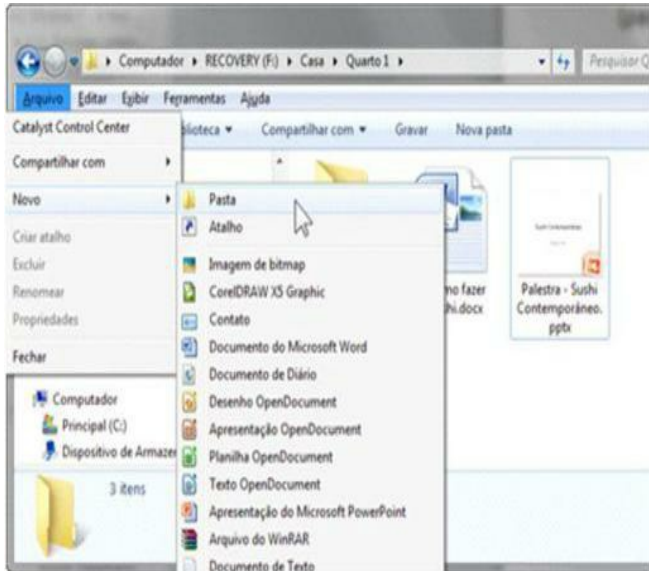
#### 4.4.1.2. Usando o Windows Explorer

Depois de conhecer os principais tópicos da interface do Windows Explorer, devemos aprender a trabalhar com ele, realizando algumas operações básicas com pastas e arquivos, como criar, renomear, excluir, copiar e mover.

A seguir, as principais operações que podemos realizar com o auxílio do Windows Explorer:

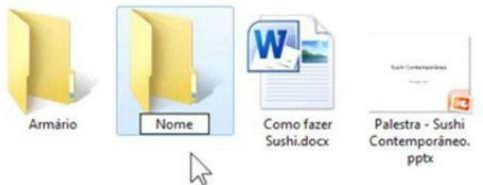
##### **Criando uma pasta ou arquivo**

Para criar uma pasta ou um arquivo, primeiro certifique-se de estar explorando a pasta ou unidade onde quer que o objeto seja criado. Acione o menu **Arquivo** (para que ele apareça, é necessário pressionar a **tecla ALT**), e, dentro dele, acione o submenu **Novo** e, por fim, clique no nome do objeto que deseja criar (em **Pasta**, por exemplo, ou no tipo de arquivo desejado).



*Figura 4.67 – Criando uma pasta.*

Após a seleção do tipo de objeto, o novo objeto será criado na pasta local, mas ele ainda precisa de um nome; basta digitá-lo (e, lógico, pressionar ENTER) e o objeto terá sido confirmado.



*Figura 4.68 – Confirmando a criação da pasta (ENTER depois de digitar o nome!).*

Esse procedimento tanto serve para pastas (conforme mostrado) como para arquivos (no submenu Novo há vários tipos de arquivos disponíveis para criar, como pode ser visto na Figura 4.67).

Outra maneira de criar uma pasta é usando o botão direito do mouse numa área em branco do Painel de Conteúdo: o **menu de contexto** vai se abrir (aliás, é o que sempre acontece quando clicamos com o botão direito do mouse em alguma coisa) e, nele, haverá a opção Novo, que é réplica do submenu Novo lá do menu Arquivo.

Lembre-se disto: sempre haverá uma forma de fazer operações no Windows 7 com o uso do botão direito (também chamado de botão auxiliar, ou botão secundário) do mouse. Ou seja, em qualquer comando aqui mostrado, sempre haverá um “jeitinho” de fazer com o botão direito! Não o subestime!

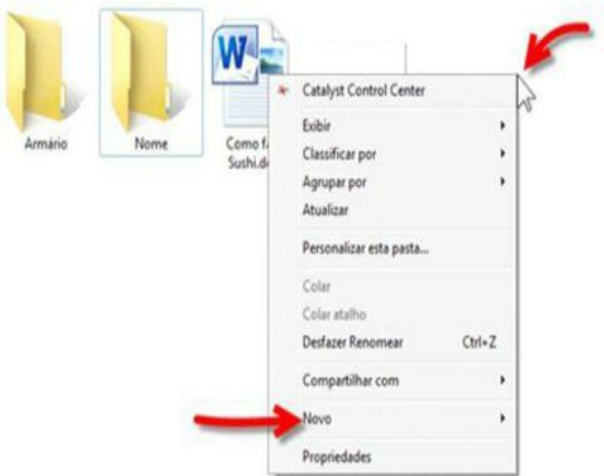
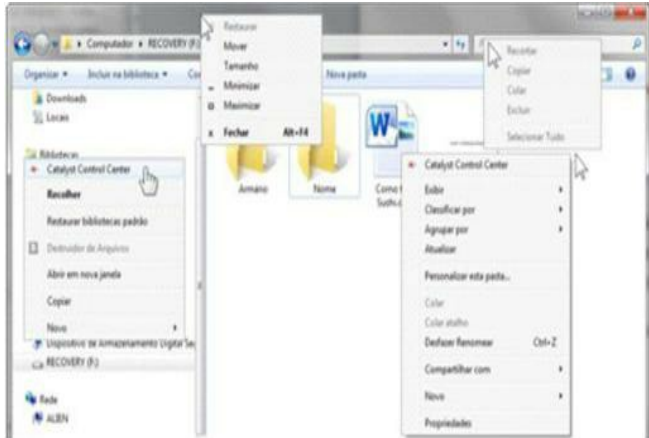


Figura 4.69 – Menu de contexto (aberto por meio do botão direito).

Portanto, para criar uma pasta, faz-se: clicar com o botão direito (numa área vazia do painel de conteúdo); clicar no submenu Novo; clicar em Pasta... Depois é só digitar o nome e ENTER para confirmar! ;-)

O menu que se abre em decorrência do clique com o botão direito é chamado de menu de contexto porque ele se adapta ao contexto! Ou seja, ele muda suas opções (comandos apresentados) de acordo com o local onde é clicado!

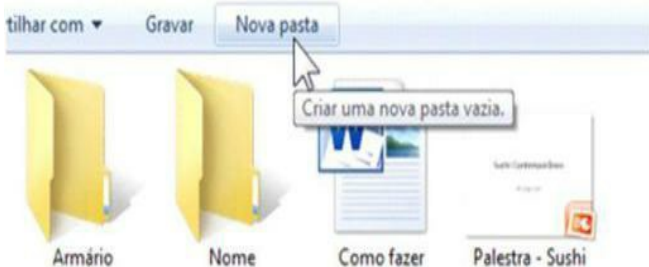
Portanto, o segredo do botão direito não é como usá-lo (é só clicar), e sim **ONDE USÁ-LO** (onde clicar).



*Figura 4.70 – Vários menus de contexto diferentes (locais diferentes).*

No Windows 7, porém, há uma forma muito mais simples de criar pastas (só pastas, ao contrário do que se viu, em que podíamos criar arquivos também, por meio das opções do submenu Novo). Basta acionar o botão Nova Pasta, na barra de ferramentas, e seguir com a digitação do nome e a confirmação! Fácil, né?



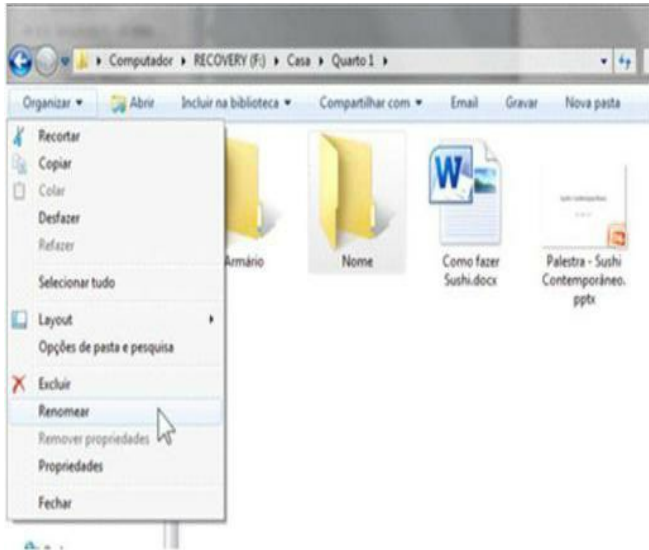


*Figura 4.71 – Taí, ó! “Facim, facim”!*

### Renomeando um arquivo ou pasta

Renomear um objeto significa mudar o nome previamente definido para ele. Para mudar o nome de um arquivo (ou pasta), siga estes passos:

1. Selecione o objeto desejado.
2. Acione o comando para renomear (há cinco maneiras):
  - a. Acione o menu Arquivo/Renomear;
  - b. Acione a tecla F2 (no teclado);
  - c. Acione um clique no **nome** do objeto;
  - d. Botão direito (no objeto)/Renomear (no menu de contexto);
  - e. Clique no botão **Organizar** (barra de ferramentas)/Renomear;
3. Digite o novo nome para o objeto (pois no nome estará alterável);
4. Confirme (pressionando ENTER ou clicando fora do objeto).



*Figura 4.72 – Usando o botão Organizar para renomear.*

Note uma coisa: no item “c” acima listado, diz-se “um clique no Nome”. Sim! É um clique só! E tem que ser no nome (não no ícone). Faça o teste!

### **Excluindo um arquivo ou pasta**

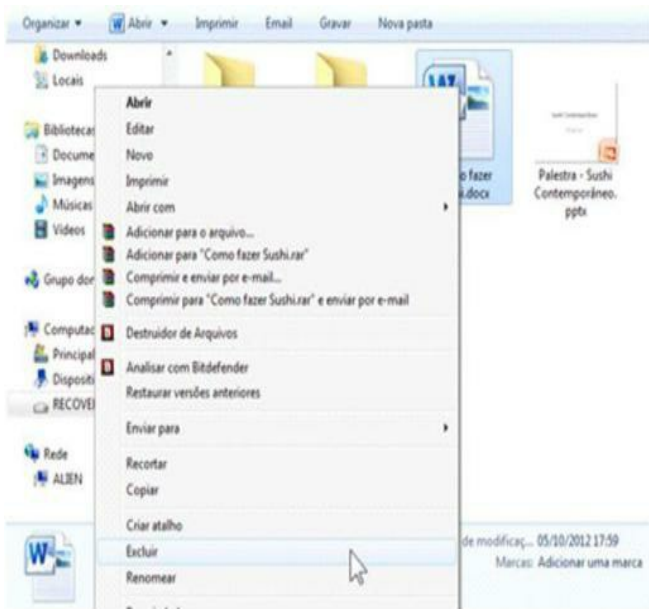
Excluir um arquivo ou pasta significa retirar este objeto da unidade de armazenamento, liberando o espaço ocupado por ele para poder ser usado na gravação de outro.

Ou seja, é “matar” o objeto, seja ele um arquivo ou uma pasta! Só lembre-se de que apagar uma pasta significa, por definição, apagar todo o seu conteúdo (todas as pastas e arquivos dentro dela).

A seguir temos um passo a passo simples para apagamento (exclusão) de um objeto:

1. Selecione o objeto desejado (ou, no caso, indesejado);
2. Acione o comando de exclusão (há quatro maneiras de acioná-lo):
  - a. Acione o menu Arquivo/Excluir;

- b. Pressione a tecla Delete (no teclado, claro!);
- c. Acione a opção Excluir do menu de contexto (clcando com o botão direito do mouse sobre o objeto a ser apagado, claro!);
- d. Acione o botão Organizar/Excluir;



*Figura 4.73 – Acionando o comando Excluir do botão direito.*

3. Confirme a operação (uma pergunta será feita em uma caixa de diálogo e toma-se por confirmação a resposta afirmativa).

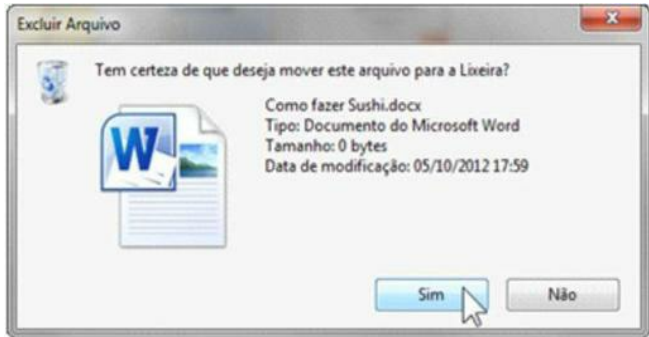


Figura 4.74 – Solicitação de confirmação de envio de um arquivo para a lixeira.

Em primeiro lugar, meu amigo leitor (ou amiga leitora), **Enviar para a Lixeira NÃO É Excluir!** Isso é uma coisa que precisa ser bem explicada! Ou seja, mesmo que a resposta à pergunta acima mostrada seja “SIM”, o arquivo em questão (Como fazer Sushi.docx) não será excluído, e sim, enviado para a Lixeira.

Outra forma de enviar um arquivo para a lixeira é arrastá-lo diretamente para o ícone da lixeira, no painel de Navegação (ou para o ícone da Lixeira na área de trabalho).



Figura 4.75 – Arquivo arrastado diretamente para a lixeira.

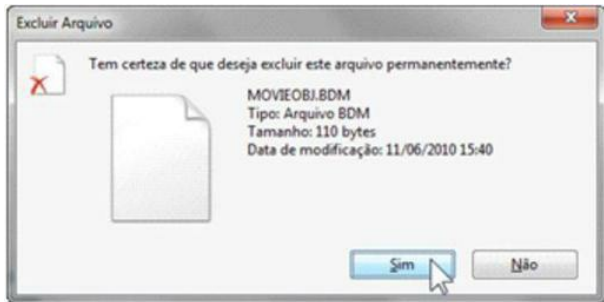
### Mas, Afinal, o que é a Lixeira?

A Lixeira é uma *pasta especial* que o sistema Windows utiliza para o processo de exclusão de arquivos e pastas dos discos rígidos do computador. A lixeira, em suma, serve para guardar arquivos que a gente tenta apagar!

Mas a lixeira só guarda arquivos que estavam em discos rígidos (HDs) ou discos de estado sólido (SSDs), que funcionam como HDs. Não importando se são discos rígidos internos ou externos (HD externo, transportável, conectado pela porta USB, por exemplo). Qualquer arquivo apagado de qualquer um desses dispositivos será, prioritariamente, armazenado na lixeira quando se tentar apagá-lo.

Arquivos armazenados em outros tipos de mídias (unidades) removíveis (como pen drives, disquetes – se ainda houver – ou cartões de memória, por exemplo) não têm direito de ir para a lixeira, ou seja, são imediatamente apagados (definitivamente).

Olha o que acontece com um arquivo armazenado num pen drive quando se tenta apagá-lo (compare a mensagem da imagem seguinte com a mensagem apresentada antes, na Figura 4.74):



*Figura 4.76 – Uia! Que medo! Agora é sério!*

Algumas “verdades e mitos” sobre a lixeira:

- a. A lixeira tem um tamanho máximo (definido pelo sistema, mas pode ser alterado por você, usuário). Sempre que a lixeira estiver cheia (de arquivos), atingindo seu “tamanho” predefinido, ela não aceitará mais arquivos.
- b. A lixeira mantém os arquivos armazenados nela por tempo indeterminado! Ou seja, nada de dizer por aí que “a lixeira apaga arquivos automaticamente depois de três dias”! O que você manda para a lixeira fica lá até que você apague de vez (esvaziando a lixeira, por exemplo) ou quando você recupera o arquivo (restaura-o para seu local original ou para outra pasta).
- c. Cada unidade de disco rígido (HD) (inclua SSD nisso, ok?) tem necessariamente sua própria lixeira. Ou seja, se um computador tem mais de uma unidade de disco rígido reconhecida (podem ser partições no mesmo disco – já que cada uma delas vai ser entendida como uma unidade diferente), cada uma delas (unidades) vai ter sua própria lixeira.

Sim: estou falando de uma pasta diferente em cada unidade de disco rígido! Essas pastas são, normalmente, invisíveis, restando, apenas, visível, a pasta Lixeira no Desktop (Área de Trabalho). Essa “lixeira central” consolida os conteúdos de todas as “lixeiras” das Unidades!

## • Unidades de Disco Rígido (2)



Principal (C:)



RECOVERY (F:)

## • Dispositivos com Armazenamento Removível (2)



Unidade de  
BD-ROM (D:)



Dispositivo de  
Armazenamento  
Digital Seguro  
(E:)

Figura 4.77 – O computador tem duas unidades de Disco Rígido reconhecidas.

No exemplo do computador acima, as unidades C: e F: possuem, cada uma delas, sua própria lixeira. São duas pastas distintas, cada uma em sua própria unidade, que são vistas juntas na Lixeira principal, lá no Desktop.

d. Os objetos presentes na lixeira (dentro dela) não podem ser abertos. (Experimente dar duplo clique em qualquer um deles! Não abre!!!).

e. Os objetos presentes na lixeira podem ser recuperados ou excluídos definitivamente. Quando se recupera um deles, ele sai da lixeira e vai para alguma outra pasta (volta a “conviver” com os demais).

Quando ele é apagado definitivamente, para o Windows, não tem mais volta!

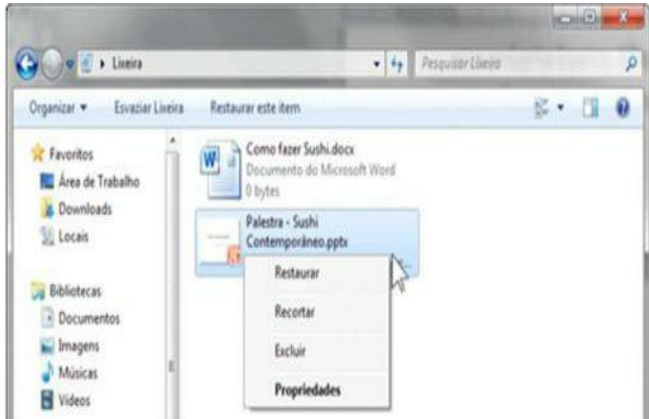


Figura 4.78 – O que se pode fazer...

Vamos analisar alguns dos comandos da lixeira (que podem ser encontrados no menu Arquivo, no botão Organizar, na Barra de Ferramentas ou por meio do botão direito do mouse):

- **Esvaziar Lixeira:** apaga, definitivamente, todos os objetos existentes na Lixeira, ou seja, “caixão e vela preta” (termo normalmente usado por mim para significar NÃO TEM MAIS JEITO!);
- **Excluir:** apaga, definitivamente, apenas o arquivo selecionado;
- **Restaurar este item (ou “Restaurar”):** envia o arquivo selecionado de volta para a pasta de onde ele foi apagado (se esta já foi apagada, ela é recriada);
- **Restaurar todos os itens:** envia todos os arquivos da lixeira de volta para seus locais originais (pastas de onde haviam sido apagados).

Um objeto que está na lixeira também pode ser arrastado para qualquer outro local fora da lixeira, sem necessariamente ir para o local de onde foi apagado.

E, ainda sobre as “verdades e mitos” da lixeira...

f. É possível ignorar a lixeira! Sim! É possível abdicar do direito de enviar um objeto para a lixeira!

Faça o seguinte: realize o procedimento de apagamento já apresentado segurando, simultaneamente ao comando, a tecla SHIFT.

(Por exemplo: SHIFT + DELETE, ou SHIFT + Arquivo/Excluir, ou SHIFT +



Organizar/Excluir... Etc.)

Quando você acionar o comando (quatro formas vistas) segurando a tecla SHIFT, o arquivo em questão, mesmo tendo direito de ir para a lixeira, será convidado a ser apagado definitivamente (ou seja, a mensagem que aparecerá diz claramente “deseja excluir o arquivo permanentemente?”).

### O Apagamento Definitivo é Mesmo Definitivo?

Bem, para começar, sabemos que os arquivos que foram enviados para a lixeira podem ser recuperados, não é mesmo?

“Sim, João, deu pra perceber!”

Mas se a pergunta fosse: “Arquivos apagados definitivamente (por exemplo, de pen drives) podem ser recuperados?”

“E aí, João? O que eu respondo?”

A resposta, caro leitor, é **DEPENDE!**

O Windows, em si, não consegue reconhecer a existência de arquivos que foram apagados definitivamente. Ou seja, o Windows não consegue recuperá-los sozinho (fazendo uso apenas de seus próprios meios e programas).

Mas há programas especiais que conseguem recuperar arquivos apagados definitivamente, mesmo de pen drives, disquetes, cartões de memória, HDs e SSDs! Tais programas podem ser achados na própria Internet... Muitos deles acompanham conjuntos de programas de segurança (como antivírus e firewalls).

“Mas, perai, João! Se os arquivos são recuperáveis, é sinal de que eles não foram apagados! Como é possível?”

Fácil, caro leitor! Vimos que os arquivos são armazenados em áreas chamadas clusters (ou unidades de alocação), nas memórias permanentes (os discos). Vimos também que esses clusters são gerenciados (organizados, controlados) por um índice normalmente chamado tabela de alocação (no Windows, é comum chamá-lo de FAT).

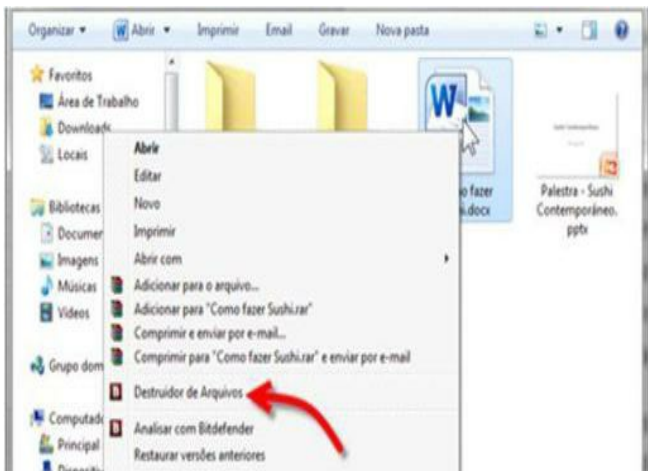
Qualquer que seja a forma escolhida para se excluir definitivamente um arquivo ou pasta do seu computador, não importando se é do disco rígido, do disquete ou de pen drives, ela apenas afetará a FAT; ou seja, um arquivo apagado ainda manterá seus dados nos clusters do disco, mas para o sistema operacional ele não existe porque a FAT informa que aqueles clusters estão vazios.

É como se, em vez de destruir uma “plantação”, o sujeito destrói apenas a “escritura da terra” – a plantação está lá... Intacta! Mas o terreno não tem mais dono! Poderá ser usado a qualquer momento para “reforma agrária”. Em suma, apagar arquivos não é destruí-los... é desapropriá-los!

Como já foi dito, há programas que conseguem ler os clusters diretamente à procura de arquivos supostamente apagados e, com isso, informar novamente à FAT sobre a presença dos mesmos (e, com isso, fazer o sistema operacional enxergá-los novamente). Esses programas são, por exemplo, usados pela Polícia Federal e Secretarias da Fazenda para vasculhar informações em computadores suspeitos de onde dados foram apagados.

Há, claro, também, programas que prometem DESTRUIR de verdade os arquivos (fazendo o que o Windows não faz), ou seja, “queimar a plantação”. Tais programas fazem o trabalho de

destruir os dados nos clusters, sobrescrevendo-os com dados aleatórios, tornando bem mais difícil (promete-se impossível) recuperar tais dados, mesmo pelos programas de recuperação.



*Figura 4.79 – Opção “Destruidor de Arquivos” – aí é matar de vez!*

A opção “Destruidor de Arquivos”, mostrada na figura anterior, vem junto com o programa BitDefender Total Security®, um conjunto de programas para segurança, que envolve antivírus, firewall e uma coleção de outros programas interessantes! Usando essa opção, é pouco provável que qualquer programa (mesmo os mais especiais) consiga recuperar o objeto “destruído”.

“Tá querendo esconder o que, João? hein? Desembucha!”

Bom... Vamos prosseguir... ;-D

Atenção: enquanto a tabela de alocação considerar que os clusters de um arquivo estão vazios (mesmo havendo ainda dados neles), eles serão considerados utilizáveis. Aí está o problema! Se algum novo arquivo for gravado naquela área em que havia dados do arquivo anterior, a recuperação do arquivo anterior fica comprometida (talvez até impossibilitada) – é o mesmo princípio do programa “destruidor”, só que sem querer!

Aviso: se você não faz ideia do que sejam clusters ou tabela de alocação, leia o capítulo anterior (a Parte 3 deste livro), que fala justamente desse assunto.

## Copiando e Movendo Objetos

É possível, através do Windows Explorer, alterar a posição de um arquivo de uma determinada pasta para outra ou criar cópias de um determinado arquivo ou pasta em outros locais.

**Mover** significa mudar um objeto de local, tirando-o do local original onde se encontra e posicionando-o em outro local (pasta). **Copiar**, por sua vez, é o procedimento que cria uma cópia exata de um determinado objeto em outro local (ou no mesmo local, desde que com outro nome).

É possível mover e copiar arquivos e pastas usando, simplesmente, o movimento de arrasto do mouse, olha só:

Para copiar um arquivo: arraste o arquivo, de seu local de origem para a pasta de destino, enquanto pressiona a tecla CTRL no teclado.

Para mover um arquivo: arraste o arquivo, de seu local original para a pasta onde deve ser colocado, enquanto pressiona a tecla SHIFT, no teclado.

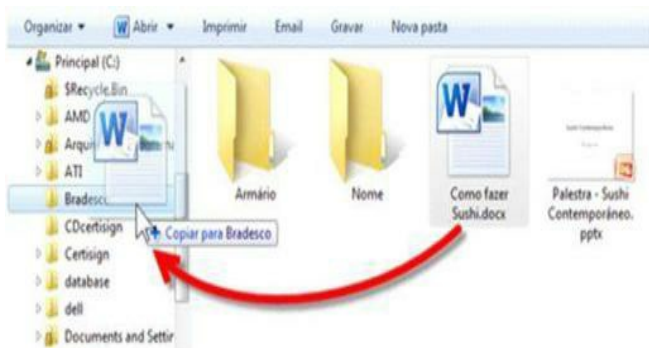


Figura 4.80 – Ao arrastar o arquivo com a tecla CTRL pressionada, o arquivo é copiado (note o indicador junto ao ícone arrastado).



Figura 4.81 – Arrastando com a tecla SHIFT pressionada, o objeto será movido.

Então, é hora de você perguntar, com a testa enrugada. “Ei, João! Nunca segurei a tecla SHIFT para mover os arquivos! Sempre arrasto os arquivos simplesmente sem segurar tecla alguma. O que é isso?” – É simples, caro leitor!

Quando o arrasto é feito *sem que se mantenha pressionada nenhuma tecla*, ou seja, quando fazemos um arrasto simples, apenas com o mouse, o resultado pode significar MOVER ou COPIAR, dependendo da seguinte condição:

a. Se o arrasto for realizado entre *pastas dentro da mesma unidade de disco*, por exemplo, se a pasta de origem e a pasta de destino do arrasto estiverem, ambas, dentro da unidade C:, então, a operação será **MOVER**.

ou

b. Se o arrasto for realizado entre *pastas de unidades de disco diferentes*, por exemplo, se a pasta de origem estiver na unidade D: e a pasta de destino do arrasto estiver dentro da unidade C:, então, a operação será **COPIAR**.

No exemplo da figura a seguir, o processo realizado é um arrasto simples. Ele será equivalente a uma cópia porque a pasta de origem (onde o arquivo está atualmente) fica na Unidade C:, e a pasta de destino (para onde o arrasto foi feito) fica na Unidade F: (portanto, unidades diferentes).

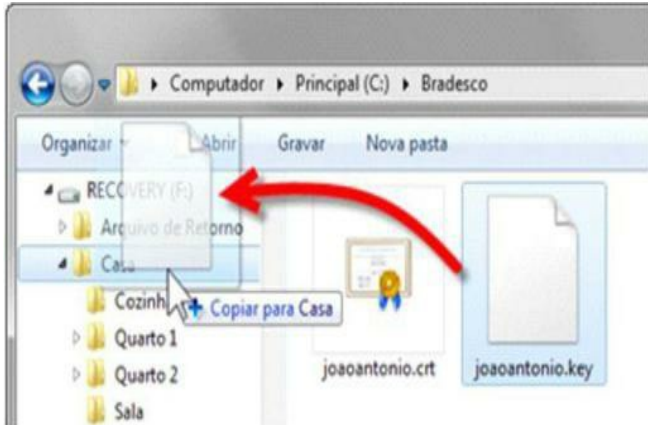


Figura 4.82 – O arrasto (sem o auxílio de tecla alguma) fazendo uma cópia.

Outra maneira de mover e copiar arquivos é usando os comandos **Recortar**, **Copiar** e **Colar**, encontrados no menu **Editar** e no botão **Organizar** (além, é claro, do menu de contexto do botão direito do mouse). Esses três comandos são usados de forma semelhante àquelas dos programas que manipulam dados, como o Word e o Excel; ou seja, os comandos Recortar e Copiar iniciam o processo, e o comando Colar SEMPRE o conclui.

Veja um passo a passo para copiar e mover arquivos usando esses comandos:

1. Selecione o objeto desejado (basta acionar um clique nele);
2. Acione o comando **Recortar** (se deseja **mover** o objeto) ou o comando **Copiar** (se deseja **copiá-lo**);
3. Selecione o local de destino (a pasta ou unidade para onde o objeto vai);
4. Acione, finalmente, o comando **Colar**.

Entenda: não importa COMO você acionou qualquer um dos três comandos (lembre-se de que pode ser pelo botão direito do mouse, pelo menu Editar ou pelo botão Organizar, na barra de ferramentas). O que importa é que você deve acionar RECORTAR ou COPIAR primeiramente (isso escolhe o tipo do procedimento que você está fazendo) e, por fim, obrigatoriamente, acionar COLAR!

Antes de você acionar o comando COLAR, nenhum procedimento foi feito! Ou seja, o procedimento só se completa quando você aciona este comando!

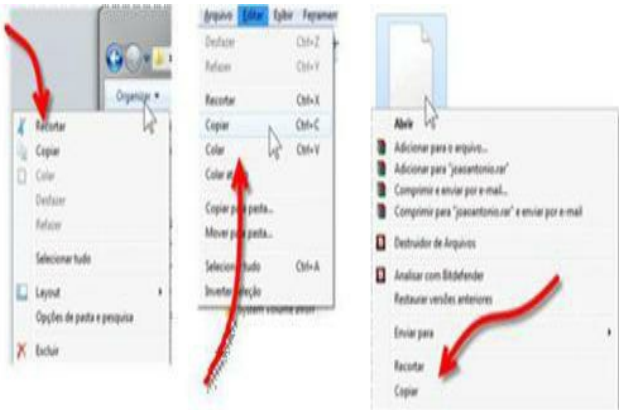


Figura 4.83 – Várias formas de acionar os três comandos!

Note, apenas, que, com relação ao botão direito do mouse, há um segredo (que, novamente, refere-se a ONDE você clica!). Se clicar num ícone de um arquivo, só aparecem as opções Recortar e Copiar (não aparece Colar). Se você clica num ícone de uma pasta ou numa área em branco (vazia) da área de conteúdo, aparece a opção Colar.

Os comandos apresentados também podem ser acionados por combinações de teclas (as chamadas teclas de atalho): CTRL + X aciona o comando Recortar; CTRL + C aciona o comando Copiar; CTRL + V aciona o comando Colar.

Novamente, vale lembrar, essas teclas de atalho são, apenas, mais uma forma de acionar os comandos! O que importa, porém, não é a forma de acionar, e sim, a sequência de acionamento.

### Comparações em Provas

Muito comum é, hoje em dia, especialmente nas provas da FCC (Fundação Carlos Chagas), que haja comparações entre os “dois métodos” de cópia e movimentação (ou seja, “arrasto” versus “três comandos”).

“Dá um exemplo, João, por favor?”

Claro!

Olha só... Se você encontra, caro leitor, a seguinte descrição numa prova: “Arrastar um arquivo de uma pasta da Unidade C: para outra pasta, dentro da unidade F:, é equivalente a acionar, depois de selecionar o referido objeto, os comandos Copiar e Colar, sequencialmente.”... O que você diria?

“Bom, João, apesar de algumas ‘estranhezas’, eu diria que está certo, porque ele comparou dois procedimentos que resultarão na cópia do arquivo!”

Precisamente!!! Ele comparou o “arrasto” entre unidades diferentes com o uso dos comandos COPIAR e COLAR, dizendo que são equivalentes! Está corretíssimo! Claro que não podemos exigir que o redator seja Ruy Barbosa (ou seja, haverá erros grosseiros de coesão, alguns até poderiam levar a interpretar a questão erroneamente!), por isso nós simplesmente abstraímos o preciosismo literário e vamos “direto na ferida”.

Ou seja, ele compara dois procedimentos e diz que são a mesma coisa (ou equivalentes)... Isso significa que ele está dizendo que os dois procedimentos dão o mesmo resultado! E... SIM! Eles dão!

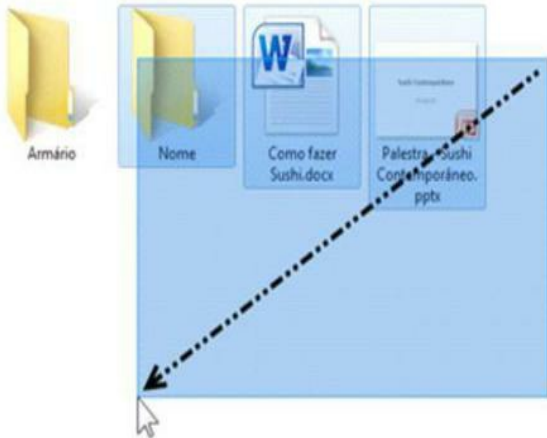
Fácil, não?!

### **Múltipla Seleção de Ícones**

Podemos realizar uma mesma operação em vários ícones ao mesmo tempo, desde que os selecionemos. Podemos selecionar ícones próximos uns dos outros (adjacentes) ou ícones que não têm contato entre si (espalhados pela janela).

As técnicas apresentadas aqui não servem apenas para o Windows Explorer, mas para todas as janelas do Windows (incluindo o Desktop). Para selecionar vários ícones próximos (adjacentes) podemos utilizar duas maneiras, a saber:

- **Quadro de seleção:** clique em uma área em branco da janela, arraste o mouse, criando um quadro, até que este envolva todos os ícones desejados. Este é o método mais fácil de fazer, mas o menos cobrado em prova (porque é difícil de “descrever” o movimento).



*Figura 4.84 – Quadro selecionando dois arquivos e uma pasta.*

- **Seleção com SHIFT:** clique no primeiro arquivo a ser selecionado da sequência e, segurando a tecla SHIFT, clique no último deles.





Figura 4.85 – Primeiro, clica-se em “Kurage sem segredos.docx” e, segurando SHIFT...



Figura 4.86 – ... Clica-se em “Sashimis Fáceis.docx” para selecionar todos entre eles.

Em tempo: Kurage (lê-se curaguê) é uma deliciosa iguaria servida em alguns restaurantes japoneses: água-viva! Sim! Água-viva! (É delicioso, apesar de nos dar a sensação de estarmos mordendo um pedaço de Tupperware® – aquelas caixas plásticas “tapauê”...)

Para selecionar vários arquivos não adjacentes (separados na tela), podemos usar a tecla CTRL.

- **Seleção com o CTRL:** clique no primeiro arquivo desejado e, segurando a tecla CTRL, clique nos demais arquivos. Pode-se, igualmente, segurar a tecla CTRL antes mesmo de selecionar o primeiro item.



*Figura 4.87 – Vários arquivos não adjacentes selecionados com a tecla CTRL.*

Para selecionar todos os ícones (arquivos e pastas) da pasta que você está explorando, é possível acionar o comando Selecionar Tudo (no menu Editar) ou acionar a tecla de atalhos CTRL + A.



*Figura 4.88 – Todos os objetos selecionados (por meio do CTRL + A)*

“João, o que eu posso fazer após selecionar vários objetos?”

Qualquer coisa, caro leitor! Apagar (excluir) todos eles de uma vez! Copiar ou Mover (arrastando ou usando os três comandos) de uma vez... Até mesmo renomear todos eles de uma única vez!

“Mas, peraí, Renomear? Eles vão ficar com o mesmo nome? Isso pode?”

Não, eles não ficarão com o mesmo nome! O Windows vai dar a eles um “(X)” no final de cada nome, onde esse “X” é um número que vai incrementando de um em um, a cada novo arquivo. Olha o resultado de ter selecionado todos os arquivos da pasta, ter acionado F2, ter digitado “Fome” e, por fim, acionado ENTER.



*Figura 4.89 – Vários arquivos renomeados ao mesmo tempo.*

Outras operações que podem ser realizadas no Windows Explorer e não envolvem arquivos ou pastas são mostradas a seguir.

### **Formatando Discos**

Formatar é preparar um disco (ou uma partição) para ser usado como superfície de gravação. Quando se formata um disco, seus dados são supostamente apagados (na verdade, a FAT tem seu conteúdo completamente apagado), deixando todos os clusters prontos para serem utilizados para a gravação de outros dados.

Lembre-se: na formatação, assim como no apagamento de arquivos, apenas a tabela de alocação é afetada. Os dados, continuam existindo em um disco recém-formatado (isso permite que sejam recuperados pelos “programas especiais” de que falei anteriormente).

Para formatar uma unidade de disco, vá ao item “Computador”, selecione a unidade a ser formatada (no nosso caso a unidade F:) e acione o comando Formatar, no menu Arquivo (lembre-se da tecla ALT).

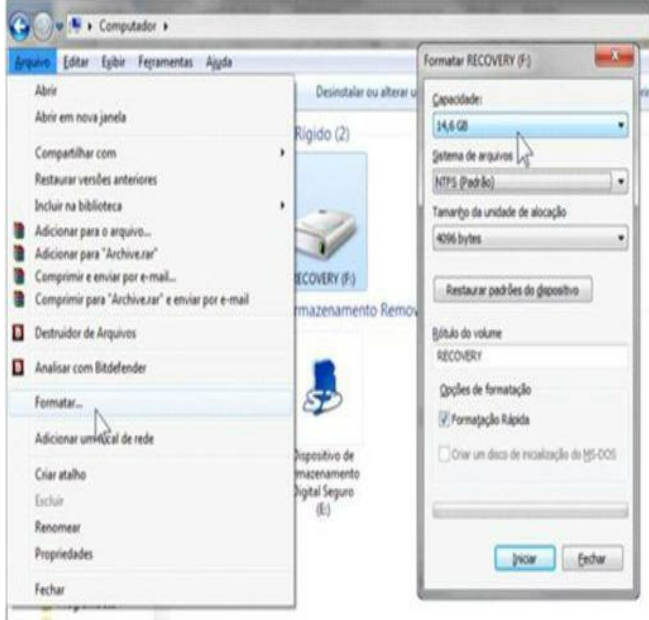


Figura 4.90 – Janela do comando Formatar aplicada a uma unidade de HD (F:).

Observe que é durante a formatação que são escolhidos o sistema de arquivos (NTFS, no nosso exemplo) e o tamanho dos clusters daquela unidade (4.096 Bytes, no exemplo).

Antes de iniciar o processo de formatação propriamente dito, é possível escolher algumas opções, a saber:

- **Formatação rápida:** que resultará, simplesmente, no apagamento da FAT (ou MFT, no caso do NTFS). Quando *não se escolhe* a formatação rápida, o Windows realiza a **formatação completa**, que significa apagar a FAT e verificar erros nos setores (clusters) após esse apagamento.

- **Rótulo do Volume:** o nome da unidade de disco (nome que é apresentado no Windows Explorer).
- **Criar disco de Inicialização do MS-DOS:** copia, para o novo disco formatado, os arquivos iniciais do sistema operacional (para que aquele disco possa ser usado para iniciar uma máquina – ou seja, fazê-la funcionar). Essa opção não está disponível para todos os tipos de discos.

Lembre-se: o Windows não deixará o usuário formatar a unidade de disco onde ele está instalado (normalmente a unidade C:). Outras unidades de disco rígido que não são importantes podem ser formatadas perfeitamente.

Lembre-se também: o comando FORMATAR pode ser encontrado, também, por meio do botão direito do mouse sendo clicado na unidade de disco que se deseja formatar!

## Compartilhando Recursos

Quando um computador faz parte de uma rede de computadores (ou seja, quando está física e logicamente conectado a outros computadores), seus recursos (unidades, pastas, impressoras) podem ser compartilhados com os outros para serem usados por qualquer componente da rede. Para compartilhar uma pasta com os outros computadores da rede, simplesmente selecione a pasta e acione **Compartilhar com...**

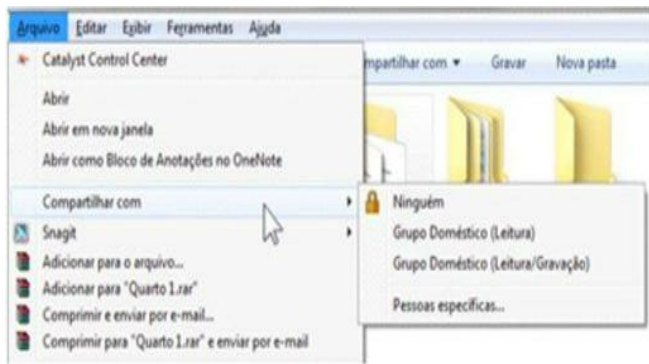


Figura 4.91 – Arquivo/Compartilhar com.

As opções que aparecem dentro de “Compartilhar com” são:

- **Ninguém:** simplesmente não compartilha a pasta selecionada. Retira todos os compartilhamentos da pasta selecionada. Depois disso, a pasta selecionada só poderá ser

acessada localmente (ou seja, do computador em que ela está) e somente pelo usuário que a criou.

- **Grupo Doméstico (Leitura):** compartilha a pasta selecionada para o Grupo Doméstico (explico depois) com o direito de “Somente Leitura” (ou seja, os demais usuários só poderão ter acesso a essa pasta para LER seu conteúdo – nunca para modificá-lo ou excluí-lo).
- **Grupo Doméstico (Leitura/Gravação):** permite compartilhar a pasta em questão para o Grupo Doméstico, dando direito, aos outros usuários, a LER e MODIFICAR (SALVAR) o conteúdo da pasta compartilhada. Ou seja, outros usuários, por meio da rede, poderão abrir (ler), salvar (gravar) e excluir os seus arquivos nessa pasta compartilhada.
- **Pessoas Específicas:** permite escolher para quem (usuários) e em quais níveis (leitura/gravação) o compartilhamento vai ser feito. Para esta opção, não é necessário ter um Grupo Doméstico.

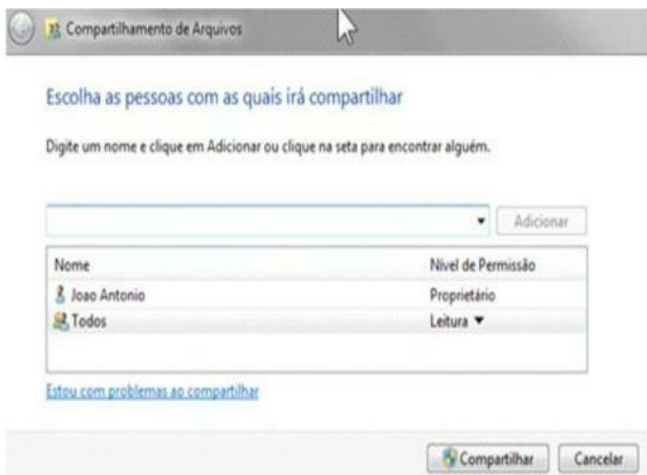
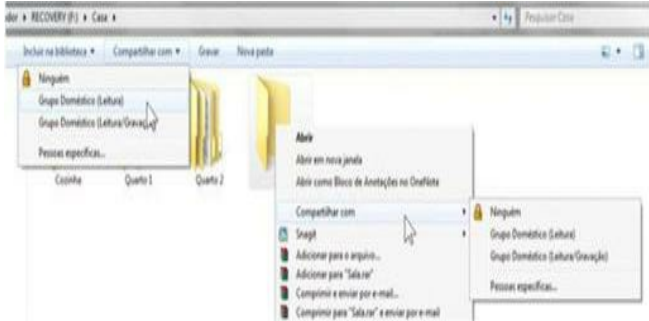


Figura 4.92 – Janela do Compartilhamento para Pessoas Específicas.

Também é possível acessar a opção **Compartilhar com** diretamente do botão direito do mouse sobre a pasta selecionada ou por meio de um botão apropriado na barra de ferramentas, conforme visto a seguir:





*Figura 4.93 – Outras formas de acionar o Compartilhar Com.*

### **Trabalhando em Rede com o Windows 7**

O Windows 7 foi criado especificamente para controlar um único computador, mas traz inúmeros recursos para que possamos trabalhar facilmente com mais de um deles ligados em rede.

Uma rede de computadores é um conjunto de computadores interligados. Uma rede permite que os vários computadores troquem informações entre si, por meio, normalmente, do compartilhamento de recursos (pastas e impressoras, por exemplo).

Por meio do Windows Explorer, é possível “ver” toda a rede. Isso é feito com a ajuda da opção Rede, que fica no Pannel de Navegação do Windows Explorer. Note o item Rede aberto, mostrando três computadores atualmente conectados (sim, eles estão ligados neste momento).

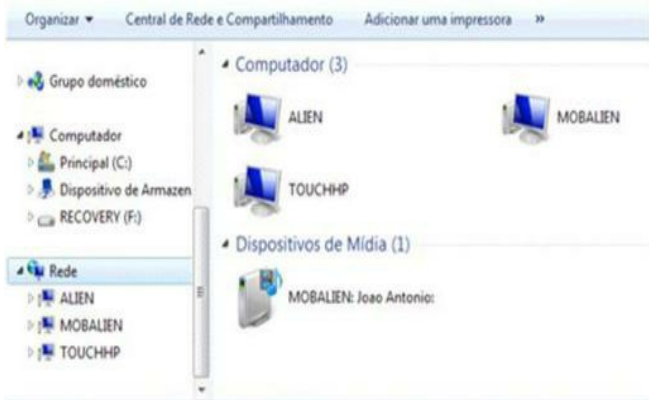
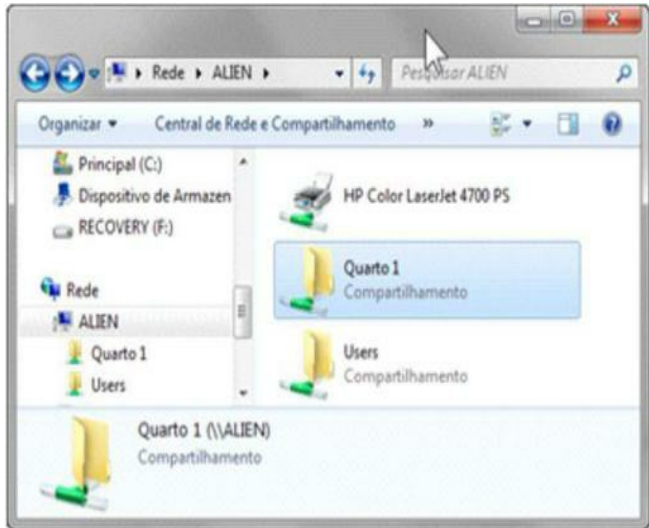


Figura 4.94 – Item Rede visualizando três computadores.

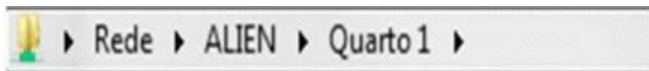
*Alien*, *MobAlien* e *TouchHP* são os nomes de três computadores ligados em rede. *MobAlien: Joao Antonio*, por sua vez, é um dispositivo de mídia, onde se poderá pesquisar músicas, vídeos e fotos – usado quando o Windows atua como Central de Mídia (Media Center) para a família.

Supondo que estamos trabalhando no computador TouchHP e queremos acessar o computador Alien, é só dar duplo clique nele! Automaticamente, seus compartilhamentos (pastas que foram compartilhadas dentro dele) vão aparecer (caso, claro, você tenha acesso, como usuário, a esse micro).



*Figura 4.95 – Em Alien, há duas pastas e uma impressora compartilhadas.*

E para entrarmos no Compartilhamento Quarto 1, basta, também, acionar um clique duplo nesta pasta. Caso tenhamos direito de acessá-la, ela será normalmente aberta! Perceba o endereço dessa pasta!



*Figura 4.96 – Endereço do Compartilhamento Quarto 1.*

Mas, se clicarmos no ícone que fica à esquerda da barra de endereços, veremos o verdadeiro endereço... Olha aí!

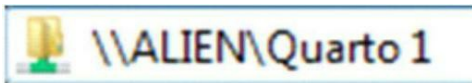


Figura 4.97 – Endereço de Rede Windows.

Explicando: quando estamos numa rede Windows (ou seja, uma rede que usa os protocolos de comunicação e compartilhamento do sistema Windows), a forma de endereçamento de outros computadores e seus compartilhamentos segue a seguinte norma:

**\\Nome\_do\_Computador\Nome\_do\_Compartilhamento**

Ou seja, sempre que nos referimos a algum computador localizado na rede, quer seja na barra de endereços, quer seja no campo de pesquisa (que fica ao lado da barra de endereços), devemos usar **\\** (*duas contrabarras*) seguido do nome do computador.

Portanto, o endereço:

**\\ALIEN\Quarto 1\Armário**

Na verdade, aponta para uma pasta chamada *Armário*, dentro de uma pasta chamada *Quarto 1*. Quarto 1, por sua vez, está compartilhada, e é localizada dentro de um computador que é conhecido, na rede, pelo nome de *ALIEN*.

Cuidado com isso, ok? Não usamos / (barra normal), e sim \ (contrabarra), exatamente como usamos nos endereços de pastas locais (do tipo C:\casa\sala).

### Mapeando uma Unidade de Rede

Mapear uma unidade de rede é selecionar um compartilhamento qualquer de outra máquina da rede e transformá-la em uma unidade de disco virtual em nosso computador. Em outras palavras, é criar um “atalho”, na forma de uma unidade de disco, que aponta para um compartilhamento em outro micro.

Para mapear uma unidade, basta acionar o menu (abra com a tecla ALT) *Ferramentas* e, lá dentro, acionar a opção *Mapear Unidade de Rede*. Dentro da caixa de diálogo que se abrirá deve-se informar a letra que a unidade usará (X:, Z:, qualquer uma) e para qual compartilhamento ela apontará.

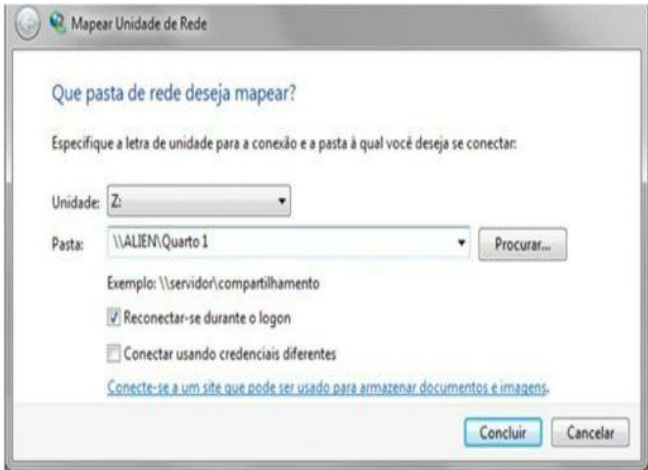


Figura 4.98 – Criando a unidade (Z:), que aponta para o “\\ALIEN\\Quarto 1”.

A opção **Reconectar-se durante o Logon** permite garantir que quando o computador for novamente ligado (e quando aquele usuário voltar a logar-se na máquina, informando suas credenciais), a unidade Z: seja novamente conectada ao compartilhamento em questão (sem precisar fazer de novo o comando Mapear Unidade de Rede).

A opção **Conectar usando credenciais diferentes** permite que se escolha um nome de usuário (login) e uma senha diferentes dos atuais (ou seja, diferentes dos usados pelo usuário atualmente ligado).

Note, na figura seguinte, como fica uma unidade mapeada, apresentada junto com as demais unidades de disco locais (veja que ela fica “separada” dos grupos “Unidades de Disco Rígidos” e “Dispositivos com Armazenamento Removível”, em um grupo próprio):

• Unidades de Disco Rígido (2)



Principal (C:)



RECOVERY (F:)

• Dispositivos com Armazenamento Removível (2)



Unidade de  
BD-ROM (D:)



Dispositivo de  
Armazenamento  
Digital Seguro  
(E:)

• Local da rede (1)



Quarto 1  
(\\ALBEN) (Z:)



Figura 4.99 – Unidade Z: – na verdade, um atalho para uma pasta na rede.

Outra forma fácil de acionar o comando **Mapear Unidade de Rede** é por meio do botão (com esse nome) na Barra de Ferramentas:



Figura 4.100 – Botão Mapear Unidade de Rede.

A qualquer momento o usuário poderá “excluir” a unidade mapeada se não a quiser mais.

Esse processo é chamado **Desconectar Unidade da Rede**, e seu comando também está localizado no menu Ferramentas.

Você também pode desconectar essa Unidade de seu compartilhamento na rede usando o botão direito do mouse sobre ela e escolhendo a opção **Desconectar-se**, no menu de contexto.

### Digitando Endereços Web no Windows Explorer

O Programa Windows Explorer é o gerenciador de arquivos do sistema operacional Windows, ou seja, tem a função de permitir a visualização e o controle dos recursos presentes nas unidades de disco do computador.

Normalmente digitamos, na barra de endereços do Windows Explorer, apenas endereços de unidades e pastas em nosso micro (ou, no máximo, endereços de pastas compartilhadas na mesma rede).

Mas é possível, sim, digitar endereços de sites e recursos disponíveis na Web (WWW). Ao teclar ENTER, depois de ter digitado o endereço da web na barra de endereços do Windows Explorer, será aberta uma janela do navegador atual (o navegador padrão configurado em seu computador) automaticamente abrindo aquele endereço digitado.



Figura 4.101 – Basta digitar o endereço da Web e acionar ENTER no Windows Explorer...



Figura 4.102 – ... E a janela do navegador padrão será aberta abrindo aquele site!

### Considerações sobre o Windows Explorer

Caro amigo leitor (ou leitora), sem dúvida alguma, o histórico das provas de concursos públicos que pediram Windows é recheado de questões de Windows Explorer (copiar, mover, excluir etc.).

Cerca de 80% de todas as questões de Windows já cobradas envolvem esse tema! Note que eu não estou dizendo que este é o único assunto que deve ser estudado por você... Longe de mim! Mas este é o mais importante!

Se você está com o tempo apertado e precisa estudar outros conteúdos e até mesmo outras matérias, limite o estudo do Windows ao Windows Explorer, está bem? É uma “dica de amigo”! Já a dica de professor é: “estude tudo!”.

#### 4.4.1.3. Bibliotecas

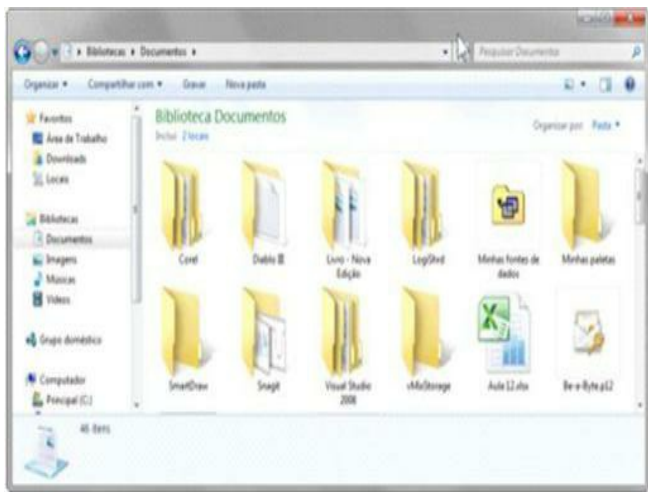
O Windows 7 trouxe um novo recurso para os usuários de computadores pessoais, recurso esse que se assemelha, um pouco, às pastas.

Uma biblioteca não é uma pasta, mas “parece”. Uma biblioteca é um conjunto (grupo) de arquivos e pastas previamente escolhidos pelo usuário. Esses arquivos e pastas podem estar em



diversos locais físicos distintos (num HD local, num pen drive, num disco compartilhado em outro micro pela rede etc.).

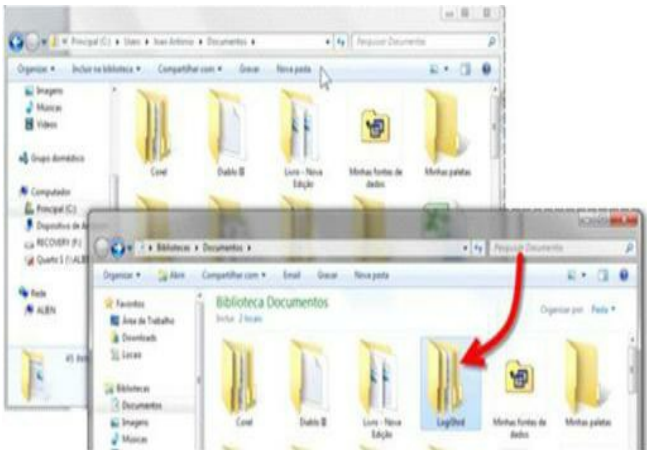
Quando abrimos uma biblioteca, visualizamos algum conteúdo nela, mas, na verdade, não temos certeza acerca de onde (fisicamente) este conteúdo está! Ele pode estar todo dentro do HD da máquina, bem como pode estar todo fora do computador (em computadores diferentes na rede)... Quem sabe?



*Figura 4.103 – Biblioteca Documentos – conteúdo de vários locais.*

Este (acima) é o conteúdo da biblioteca Documentos. A biblioteca Documentos não é a pasta Documentos (sim, essa pasta existe, dentro da pasta pessoal do usuário). A biblioteca Documentos normalmente aponta, sim, para o conteúdo da pasta Documentos, mas também pode apontar, além dela, para outras pastas.

Bibliotecas são “formas centralizadas” de visualizar o conteúdo de várias pastas distintas, quer estejam dentro do micro, quer estejam em outros computadores acessíveis por meio da rede.



*Figura 4.104 – Pasta “LogiShrd” existe na biblioteca, mas não na pasta Documentos.*

No exemplo da figura acima, vemos o conteúdo da pasta Documentos (em cima) e o conteúdo da biblioteca Documentos (onde existe a pasta LogiShrd). Note: esta pasta não existe dentro da pasta Documentos, somente na biblioteca Documentos, provando que não são a mesma coisa.

Então, entenda que:

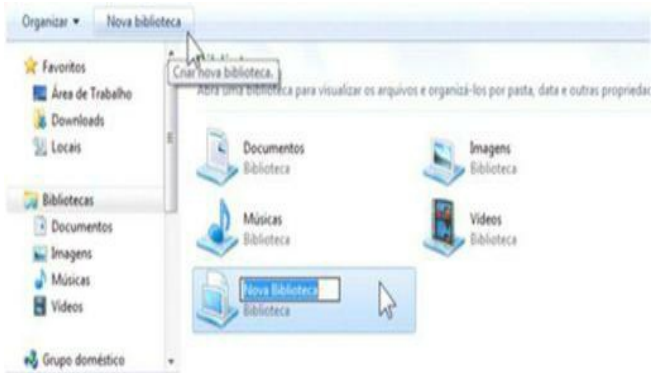
- a. Uma biblioteca não é uma pasta.
- b. Uma biblioteca pode ser aberta e visualizada como se fosse uma pasta.
- c. Existem quatro bibliotecas padrão: Documentos, Vídeos, Imagens e Músicas.
- d. Novas bibliotecas podem ser criadas.
- e. Todas as bibliotecas ficam num único local: chamado de “Bibliotecas”, acessível pelo painel de navegação do Windows Explorer.



Figura 4.105 – Windows Explorer, acessando as bibliotecas.

### Criando uma Biblioteca Nova

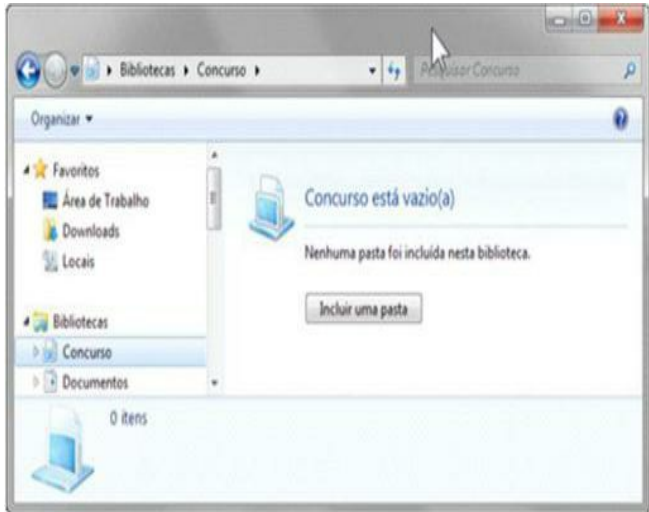
Para criar uma nova biblioteca, basta acionar o botão “**Nova Biblioteca**”, na barra de ferramentas, ou o menu (lembre-se da tecla ALT) **A**rquivo, depois acionar **N**ovo, e, em seguida, **B**iblioteca.



*Figura 4.106 – Criando uma Nova Biblioteca.*

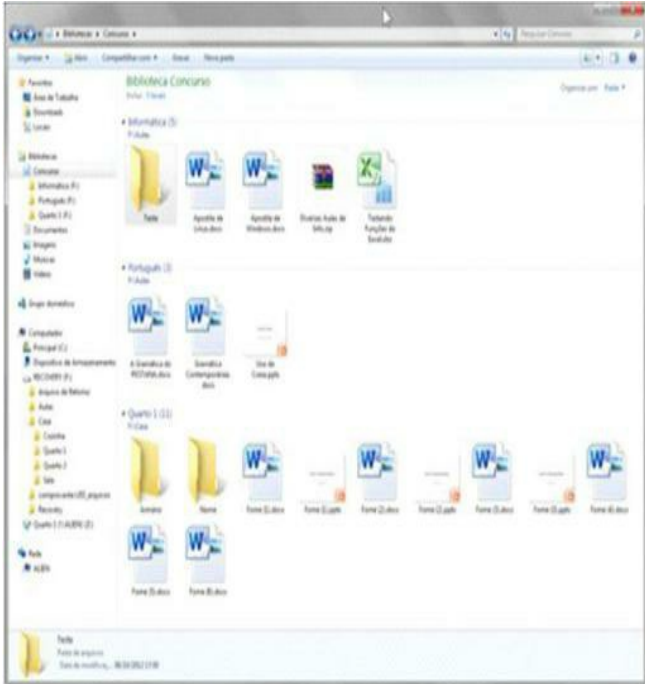
Digite o novo nome para a biblioteca e ela estará pronta! Depois disso, basta acionar duplo clique sobre ela para abri-la. Você verá que o Windows entenderá que ela está vazia, por isso, pedirá que você indique quais pastas vão ter seu conteúdo visualizado por essa biblioteca...

Não se esqueça! Bibliotecas são “formas de visualizar” o conteúdo de pastas. Você pode associar uma biblioteca a várias pastas! Desta forma, o que você vê, na biblioteca, é a união dos conteúdos de todas as pastas associadas àquela biblioteca.



*Figura 4.107 – Chamei minha biblioteca de Concurso e tentei abri-la!*

Depois de adicionar três pastas à minha biblioteca, ela passa a mostrar o conteúdo das três pastas, inicialmente separando o conteúdo pelas pastas originais. Mas a forma de visualizar pode ser mudada pelo Modo de Exibição!



*Figura 4.108 – Biblioteca Concurso apontando para três pastas.*

Quando você realiza qualquer operação dentro da janela da biblioteca, como, por exemplo, apagar um arquivo, fique ciente, caro leitor, de que ele será apagado DA PASTA DE ORIGEM. Porque, simplesmente, a biblioteca é uma “forma de visualizar”, e não uma pasta nova!

Os arquivos vistos na biblioteca não são cópias! São os próprios arquivos, existentes em suas pastas originais! Só isso! A biblioteca é um grande bisbilhoteiro, que “enxerga” os arquivos em

suas pastas originais!

#### 4.4.2. Painel de controle

O sistema operacional Windows 7 traz, como suas versões anteriores, um programa chamado **Painel de Controle**, que permite configurar com detalhes os diversos aspectos do programa.

O painel de controle é, em poucas palavras, uma janela cheia de ícones, e cada um desses ícones representa um quesito específico para ser ajustado para o Windows. Ou seja, cada item “mete seu nariz” em um aspecto diferente, permitindo o controle e a configuração total do sistema. Veja, a seguir, a janela do painel de controle em seu modo **Ícones Grandes** (neste modo, cada ícone é responsável por um aspecto de ajuste do Windows);



Figura 4.109 – Painel de Controle no Ícones Grandes.

Mas o painel de controle também pode ser apresentado de outra forma, o que, para os tradicionalistas conhecedores das versões anteriores do Windows, é um martírio: o **Modo**

**Categorias**, em que os itens são divididos em categorias, cada qual com seu nome (e ícone) específico.



*Figura 4.110 – Painel de Controle no modo de categorias.*

Não iremos analisar o painel de controle item a item neste livro porque isso o tornaria muito maior! Além disso, painel de controle não tem sido tão cobrado em prova ultimamente (lembre-se de que o Windows Explorer é mais importante!).

Mas isso não quer dizer que você não terá este conteúdo! Acesse o site da Editora Campus/Elsevier, e pegue, no hotsite da Série Concursos, a apostila de Painel de Controle (um arquivo PDF com algumas dezenas de páginas). É só para não dizer que eu não tive o trabalho de preparar isso para você, ok?

#### **4.4.3. Acessórios do Windows**

São alguns pequenos aplicativos que acompanham o sistema operacional Windows. Esses programas têm sérias limitações de uso por não serem profissionais, mas na falta de outro, eles



“quebram um galho”.

Todos esses programas podem ser achados dentro da opção *Acessórios*, que se encontra no menu *Todos os Programas*, do menu *Iniciar*.

#### 4.4.3.1. Calculadora

A Calculadora do Windows simula uma calculadora de bolso e apresenta alguns formatos (modos de exibição e funcionamento) interessantes.



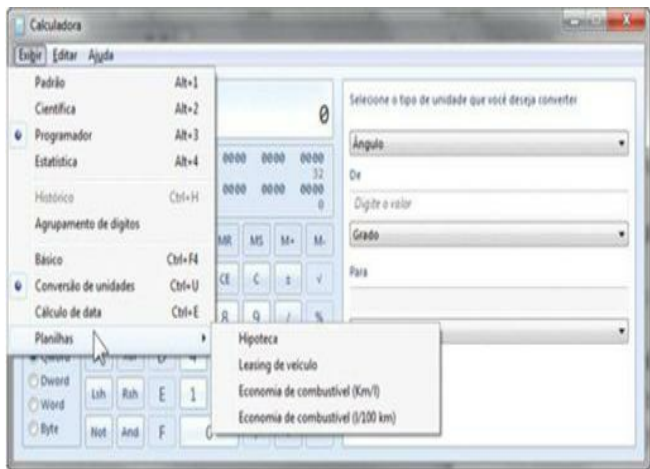
Figura 4.111 – Calculadora no formato científico.

A calculadora pode apresentar-se em quatro modos distintos:

- **Padrão:** apenas com algumas operações matemáticas básicas.
- **Científica:** contemplando mais operações, comuns às calculadoras científicas.
- **Programador:** contendo também operações de conversão de base numérica (binário, decimal, octal e hexadecimal), além de operações booleanas (como AND, OR, XOR, NOT).

- **Estatística:** contendo funções como somatórios, frequências, e outros recursos para cálculos estatísticos.

Além desses quatro formatos de funcionamento, o menu Exibir da Calculadora oferece alguns recursos especiais, como podemos ver abaixo:



*Figura 4.112 – Recursos especiais da Calculadora.*

São formulários pré-programados para cálculo de hipoteca, leasing, consumo de combustível, além de conversão de unidades de medida e cálculos com datas. É, meu amigo leitor... a Calculadora do Windows evoluiu!

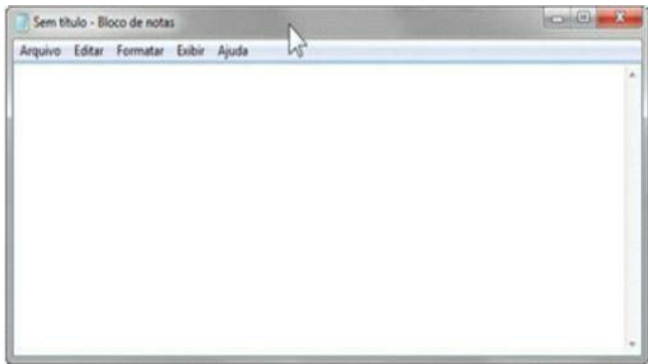
#### 4.4.3.2. Bloco de notas

Pequeno programa classificado como **editor de textos** que acompanha o Windows. O Bloco de notas é classificado como editor de textos porque permite uma forma bem simples de edição, apenas escreve e apaga caracteres puros (em código ASCII). Nesse aplicativo não há formatação (negrito, itálico, sublinhado, fontes, cores etc.) nem recursos extras (tabelas, figuras, marcadores, numeração etc.) como no Word.

“João, você está enganado! Há opções de formatação no Bloco de notas! Sempre escolho

fontes na opção Formatar/Fonte do programa.”

Sim, leitor, mas essa opção não formata o texto em si. Ela formata apenas a apresentação do texto na tela, porém o texto continua como um texto sem formatação, ou seja, sem efeitos de fonte atrelados a ele.

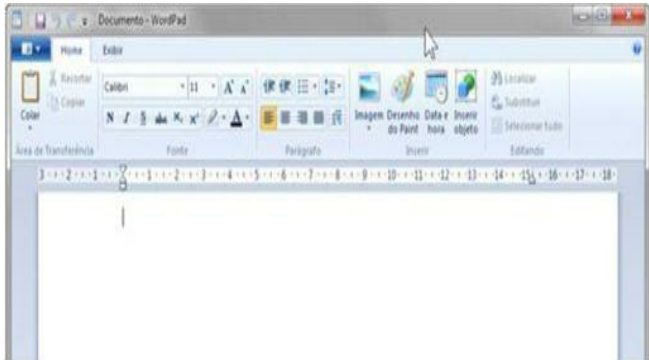


*Figura 4.113 – Bloco de notas.*

O Bloco de notas, assim como qualquer programa editor de texto, é ideal para programação. Isso porque os programas (códigos) só podem ser escritos em texto puro (texto simples).

#### **4.4.3.3. Wordpad**

Classificado como processador de textos por possuir recursos de formatação e alguns efeitos a mais, o Wordpad é, na verdade, uma versão simplificada do Microsoft Word.



*Figura 4.114 – Wordpad do Windows 7.*

No Windows 7, os arquivos do Wordpad são salvos, por padrão, no formato RTF, mas o programa também pode salvar (e abrir) arquivos TXT, DOCX e até mesmo ODT (do BrOffice).

#### **4.4.3.4. Paint**

Programa de pintura que acompanha o Windows. O Paint permite que o usuário crie e edite arquivos de bitmap (imagens formadas por pequenos pontos coloridos – os pixels). O Paint não trabalha com imagens vetoriais (desenhos feitos através de cálculos matemáticos), apenas permite a pintura de pequenos pontos para formar a imagem que se quer.



*Figura 4.115 – O Paint trabalha com bitmaps (imagens com pequenos quadradinhos).*

Os arquivos feitos pelo Paint são normalmente salvos com a extensão PNG, mas o programa também permite salvar os desenhos com outros formatos de arquivos de imagem, como JPG (JPEG), GIF, TIFF e BMP (seu “antigo formato padrão”).

#### 4.4.3.5. Outros acessórios do Windows 7

Há alguns outros programas no menu Acessórios do Windows 7 que podem ser úteis, e interessantes em provas. Vamos a eles:

- **Central de Sincronização:** permite sincronizar arquivos (versões ou cópias existentes em computadores diferentes), para manter todas as cópias atualizadas, com o mesmo conteúdo.
- **Conectar a um Projetor e Conectar a um Projetor de Rede:** são dois pequenos programas que fornecem auxílio para conectar o computador em questão diretamente (ou por meio da rede) a um projetor.
- **Conexão de Área de Trabalho Remota:** permite acessar remotamente (de longe) um computador, a fim de controlá-lo como se estivéssemos diante dele. É necessário saber o endereço do computador a ser controlado e ter a senha para realizar esta operação.
- **Notas Autoadesivas:** apresenta, na tela do Windows, pequenas janelas com o formato de adesivos (post-it) para “recadinhos”.
- **Prompt de Comando:** abre uma janela para dar comandos ao Windows em interface textual (como se fazia no DOS).
- **Painel de Entrada de Expressões Matemáticas:** permite usar uma caneta ou interface sensível ao toque (tablets, monitores touchscreen) para escrever e interpretar expressões matemáticas.

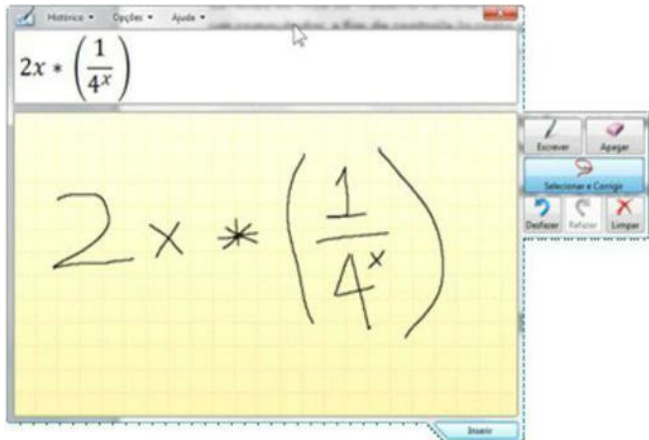


Figura 4.116 – Painel de Entrada de Expressões Matemáticas... Útil?

- **Windows Mobility Center:** permite configurar alguns ajustes interessantes para micros portáteis (PC Móvel), como laptops e notebooks. Entre as configurações, estão: brilho do monitor, consumo de energia, orientação da tela, configurações para conexão com projetor, entre outros.

#### 4.4.4. Ferramentas do sistema

Dá-se o nome de Ferramentas do sistema a um conjunto de programas utilitários que vêm junto com o Windows. Esses programas visam “consertar” certos problemas do computador, melhorando seu desempenho.

Todos eles podem ser encontrados em **Iniciar/Todos os Programas/Acessórios/Ferramentas do Sistema:**

Vamos a eles:

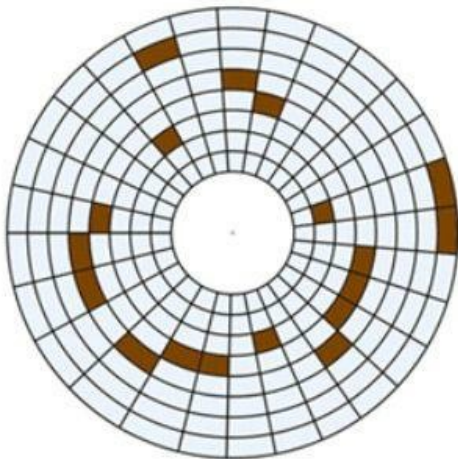
##### 4.4.4.1. Desfragmentador de disco

Ferramenta que organiza os clusters em uma unidade de disco. Pode ser que na prova eles afirmem simplesmente (de forma bem minimalista) “... O desfragmentador de disco organiza os arquivos e pastas no disco...”. A frase não está certa, porque o que é organizado é a estrutura de

clusters em si, mas é uma forma bem “superficial” de descrever sua função. (Já vi em provas essa frase ser considerada verdadeira!)

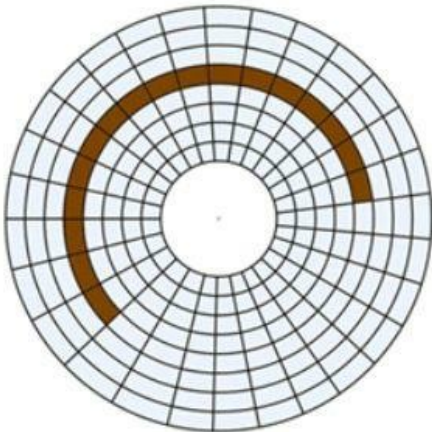
O que realmente o desfragmentador faz é reunir os clusters (blocos) que fazem parte de um mesmo arquivo para que fiquem em posições contíguas na unidade de disco, objetivando, com isso, a aceleração da leitura e gravação na referida unidade. Vamos à explicação mais detalhada.

Quando usamos um computador, há um processo natural que acontece nas unidades de disco: a fragmentação. Ou seja, os arquivos gravados em vários clusters têm seus pedaços “separados” pela superfície do disco. Isso é uma consequência natural do uso dos discos. Na figura a seguir, pode-se ver, com certo exagero, um arquivo gravado num disco. (Ele pode chegar a ser assim!)



*Figura 4.117 – Arquivo fragmentado no disco.*

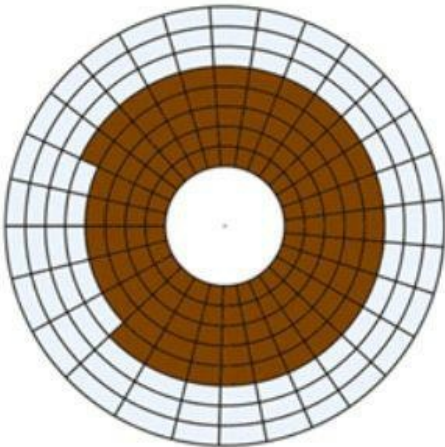
O desfragmentador une os pedaços dos arquivos de forma que os blocos do arquivo fiquem em sequência, para facilitar a leitura por parte do dispositivo mecânico que guia o braço da cabeça de leitura/gravação da unidade de disco. Depois de desfragmentar uma unidade de disco rígido, será perceptível a melhoria de seu desempenho.



*Figura 4.118 – Arquivo já desfragmentado.*

O programa desfragmentador também organiza o espaço livre na unidade, separando-a da área onde há espaço ocupado. Durante o uso normal dos discos, há intercalações entre blocos usados e blocos livres em sua superfície, devido aos diversos processos com arquivos (apagar, criar, mover, copiar etc.), e o desfragmentador separa os blocos usados (que organizadamente vão para a parte mais central – o início – do disco) e os blocos livres (que vão parar na área mais periférica – o final – do mesmo).





*Figura 4.119 – O espaço livre é separado do espaço ocupado após a desfragmentação.*

#### **4.4.4.2. Monitor de Recursos**

Apresenta, em tempo real (ou seja, constantemente) as informações acerca de uso dos recursos do computador, como CPU, memória, memória virtual, entre outros.

Clicando nos nomes dos programas em execução (programas abertos naquele momento), é possível visualizar o quanto estão ocupando e gastando dos recursos do computador nos gráficos à direita.

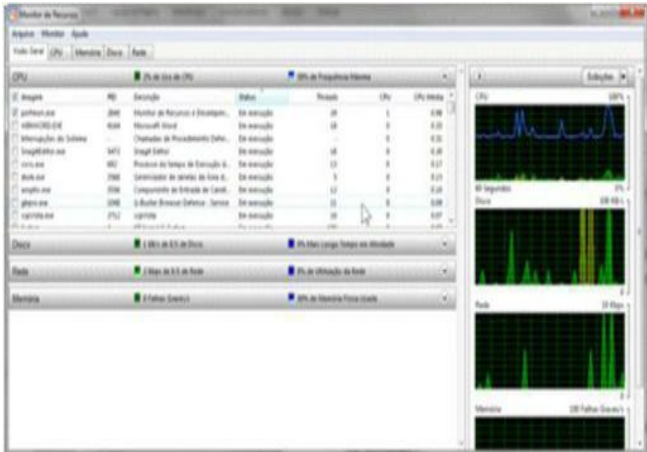
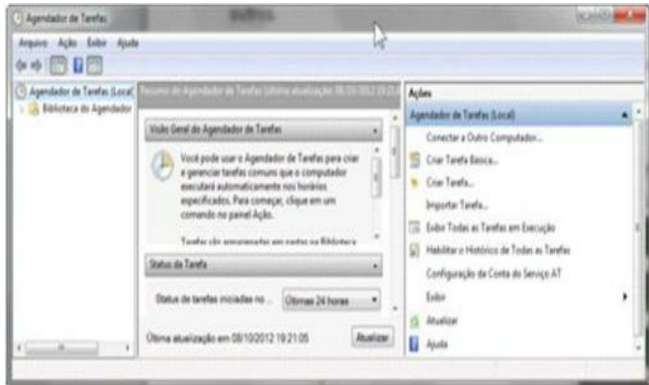


Figura 4.120 – Ferramenta Monitor de Recursos do Sistema.

#### 4.4.4.3 – Agendador de Tarefas

Este programa permite definir datas e horários específicos para a execução de certos programas em seu sistema.

Agendar para executar o antivírus, por exemplo, todas as sextas-feiras à noite é um exemplo de utilização deste programa.



*Figura 4.121 – Agendador de Tarefas.*

#### 4.4.4.4. Limpeza de disco

É um utilitário que vasculha as unidades do computador à procura de arquivos que possam ser apagados pelo usuário a fim de liberar mais espaço nesses discos.

O utilitário de Limpeza de disco sugere que podem ser apagados os arquivos que estão na lixeira (que já deveriam ter sido apagados pelo usuário), os arquivos temporários da Internet (fotos, páginas, vídeos e tudo o mais que se adquire navegando na Web) e os arquivos temporários que o sistema operacional Windows não apagou. Alguns outros arquivos que o programa julga desnecessários são apresentados na lista mostrada na figura seguinte:

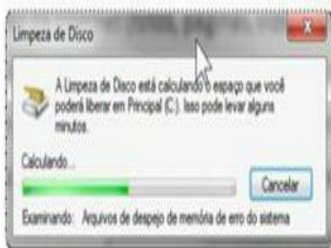


Figura 4.122 – Limpeza de disco na unidade C:.

#### 4.4.4.5. Restauração do sistema

Este recurso permite que o Windows desfça alterações realizadas pela instalação de algum programa no sistema e restaure as configurações em vigor antes desta instalação.

A restauração de sistema retorna, normalmente, ao estado do Windows quando este foi instalado no computador. Caso o usuário queira que o Windows retorne a um estado mais recente, deve criar um **Ponto de Restauração**, que seria uma descrição completa de como o Windows está naquele determinado momento. Uma vez criado o ponto de restauração, o sistema pode ser recuperado a qualquer momento e retornar àquele estado exato.

Esse utilitário é muito interessante, pois existem vírus de computador e outros programas maliciosos que adoram tirar o sono dos usuários alterando as configurações do sistema Windows.

Um exemplo bem simples: imagine que seu computador está funcionando perfeitamente, e que você decidiu instalar aquele jogo que comprou em uma banca de revistas qualquer. Claro que pode acontecer algo, não é? Para se prevenir, você usa o recurso de Restauração do Sistema a fim de criar um ponto de restauração antes de instalar o jogo.

Se o jogo criar algum problema de instabilidade no sistema (por exemplo, fazendo o Windows trabalhar muito mais lento do que trabalhava antes), você pode solicitar que o Windows retorne ao estado anterior ao momento da criação do ponto de restauração, ou seja, o seu sistema vai voltar a funcionar exatamente como estava antes da instalação do jogo. Ele simplesmente vai passar a entender que a instalação do jogo jamais aconteceu!

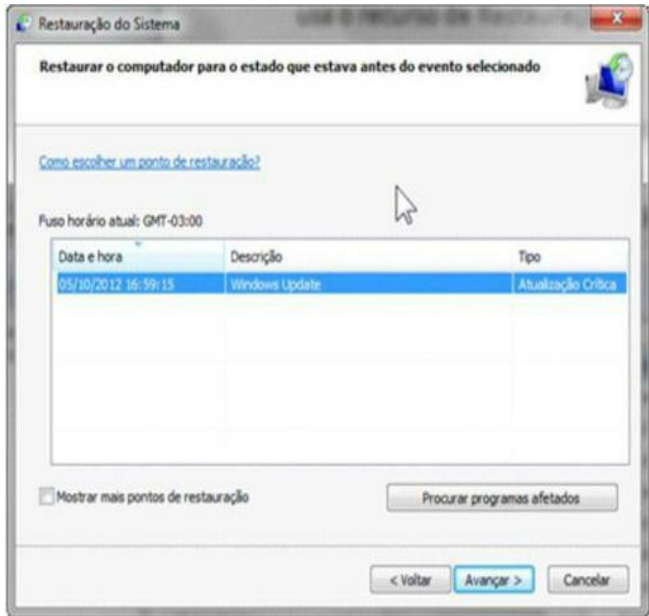


Figura 4.123 – Restauração do sistema.

No Windows 7, os itens Painel de Controle e Computador (item que apresenta as Unidades de Disco) também estão presentes dentro do grupo Ferramentas de Sistema.

#### 4.5. Outras dicas sobre o Windows

Bem, o sistema operacional Windows não é um dos assuntos preferidos em concursos públicos, mas, de vez em quando, aparece uma questão sobre ele! Há algumas outras dicas a respeito do Windows a serem mostradas neste material.



#### 4.5.1. Combinações com a tecla (Windows)

Além de CTRL, SHIFT e ALT com as quais nunca nos acostumamos, a tecla Windows (vista no início desta parte sobre Windows) pode ser usada em combinações com outras teclas para acionar comandos mais rapidamente. Conheça as combinações:

| Acione...                                                                             | Para...                                                                          |
|---------------------------------------------------------------------------------------|----------------------------------------------------------------------------------|
|  + E | Abrir o Windows Explorer                                                         |
|  + F | Abrir a janela para pesquisar arquivos e pastas                                  |
|  + R | Abrir a janela do comando Executar                                               |
|                                                                                       | Mostrar a área de trabalho (o Desktop) – esse comando tanto é usado para mostrar |



+ D

o Desktop  
(minimizando  
todas as janelas  
abertas) quanto  
para voltar as  
janelas ao seu  
estado original



+ M

Minimiza todas as  
janelas. Esse  
comando não as  
faz voltar ao  
tamanho original  
(ou seja, é um  
caminho sem  
volta)

Bloqueia a  
Estação de  
Trabalho (o



+ L

computador). Para desbloqueá-lo, o Windows solicitará a senha do usuário. Durante a tela de bloqueio, é possível usar o recurso de Trocar Usuário (visto adiante)



+ F1

Abre a janela de Ajuda e Suporte do Windows

Abre a janela de configuração rápida de conexão





+ P

com projetor, permitindo que o usuário defina se o projetor vai apresentar o mesmo conteúdo da tela principal ou não



+

TAB

Realiza o Flip 3D (alternância de janelas com estilo tridimensional)

Realiza a visualização da área de trabalho (deixa todas as janelas translúcidas, como



+

ESPAÇO

se fossem de vidro) – efeito igual a manter o ponteiro do mouse no botão Mostrar Área de Trabalho (na extremidade direita da barra de tarefas)



+ T

Alterna entre os botões abertos na Barra de Tarefas (pode-se acionar uma vez e navegar pelas setinhas do teclado, ou acionar o “T” várias vezes, com a tecla

|                                                                                       |                                                                                                                                                                                                                                |
|---------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|                                                                                       | Windows<br>acionada)                                                                                                                                                                                                           |
|  + X | Aciona o<br>Windows Mobility<br>Center (se a<br>versão possuir)                                                                                                                                                                |
|  + U | Aciona o utilitário<br>Central de<br>Facilidade de<br>Acesso, para<br>configurar itens de<br>acessibilidade de<br>usuário, como<br>Lupa, Teclado<br>Virtual, Narrador<br>(leitor de tela) e<br>Alto Contraste,<br>entre outras |

## 4.5.2. Atributos dos arquivos

Como todo sistema operacional, o Windows grava os arquivos em seu disco com algumas “características” próprias, que chamamos de **atributos**.

Quando clicamos com o botão direito do mouse em um arquivo e acionamos o comando **Propriedades**, temos acesso às informações a respeito do arquivo, como data de criação, nome, tamanho e também podemos ver seus atributos.

Além dos dois primeiros atributos apresentados na parte inferior da janela, temos acesso aos Atributos Avançados, por meio do botão **Avançados...**, também nesta janela.

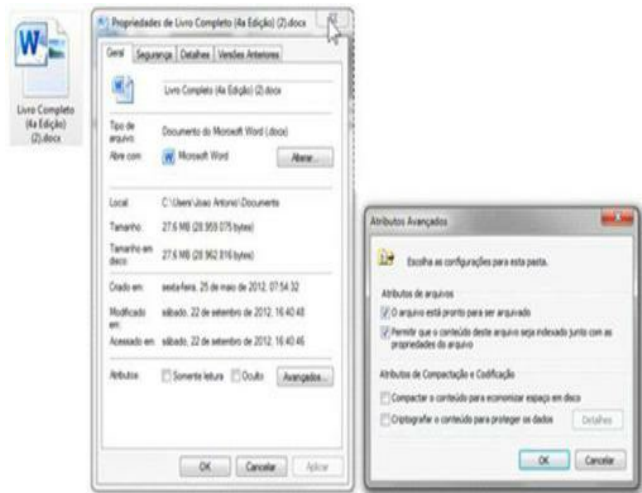


Figura 4.124 – Janela Propriedades do Arquivo e caixa Atributos Avançados.

A estrutura com a qual o Windows grava seus arquivos define três atributos possíveis a qualquer arquivo:

- **Somente Leitura:** define que o arquivo não poderá ser salvo, apenas lido. Ou seja, um arquivo marcado com esse atributo não pode ser modificado a menos que se retire a definição de Somente Leitura.
- **Oculto:** define que o arquivo não será visto nas janelas do Windows Explorer. Só é possível acessar esse arquivo se o nome dele for conhecido.

E dentro da janela Atributos Avançados (acessível por meio do botão Avançados...):

- **O arquivo está pronto para ser arquivado (antigamente chamado de “Arquivo Morto”, ou “Arquivamento”):** define que o arquivo em questão participará do próximo backup a ser realizado no computador. Esse atributo só é interessante para programas de backup.
- **Permitir que o conteúdo do arquivo seja indexado junto com as propriedades do arquivo:** inclui o conteúdo do arquivo na tabela de índice de pesquisa do Windows (essa tabela normalmente contém apenas as propriedades básicas do arquivo, como data de modificação, data de criação, tamanho etc.).
- **Compactar o conteúdo:** grava o arquivo no disco de forma compactada, assim, o arquivo é armazenado consumindo muito menos bytes em disco.
- **Criptografar o conteúdo:** grava o arquivo no disco, escrevendo-o de forma embaralhada (criptografada), assim, ele só poderá ser aberto pelo usuário que o criou, quando este faz seu login no Windows.

#### 4.5.3. Windows Update

Recurso que permite ao Windows se conectar aos servidores da Microsoft para se “atualizar” com os novos componentes e programas que a Microsoft coloca à disposição dos usuários.

Regularmente a Microsoft coloca, na Internet, pequenos programas corretivos ou atualizações do Windows para que os usuários possam ter sempre um sistema operacional novo e seguro (pelo menos, é esse o intuito). O Windows Update é uma página da Internet que faz a busca dessas novidades nos servidores da Microsoft e as instala no computador (com a autorização do usuário).

O Windows Update é encontrado em Iniciar/Todos os Programas/Windows Update, ou por meio do Painel de Controle, na categoria Sistema e Segurança.

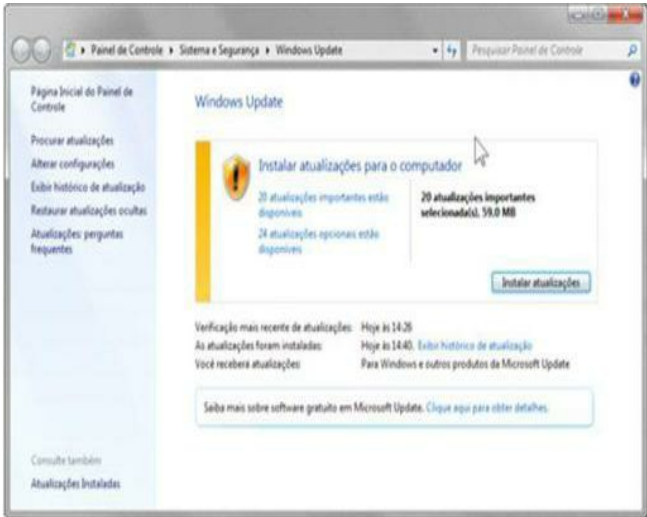


Figura 4.125 – Windows Update (é necessário atualizar, hein?).

Há três tipos de atualizações que podem ser adquiridas através do Windows Update (são classificadas por ordem de importância):

- **Atualizações Importantes:** normalmente são associadas às atualizações críticas de segurança e privacidade. Corrigindo falhas recém-descobertas. Essas correções trazem, também, melhorias na confiabilidade do sistema Windows.
- **Atualizações Recomendadas:** incluem normalmente atualizações de software e novos e/ou aperfeiçoados recursos para o computador.
- **Atualizações Opcionais:** trazem recursos extra para os programas, que podem ser instalados manualmente. Não há, porém, qualquer necessidade de instalação destes recursos. Você faz se quiser! (Um bom exemplo desse tipo de atualização são os pacotes de idiomas extra do Windows 7 Ultimate.)

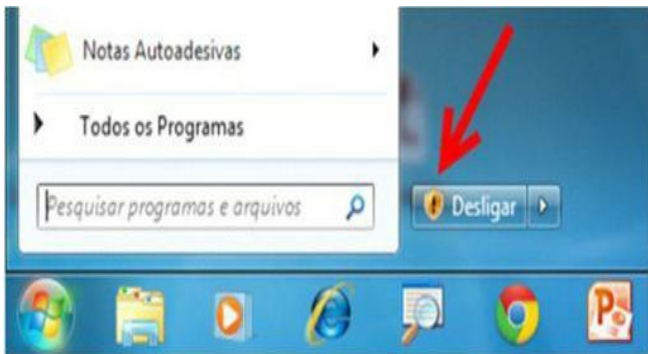
As atualizações “isoladas”, uma a uma, são, normalmente, conhecidas como **Patches** (“curativos”).

De vez em quando, a Microsoft libera, de uma vez só, um grande pacote de atualizações

reunidas, aliadas a novos recursos. Essa “operação plástica” é chamada de *Service Pack* (“Pacote de Serviços”).

Um Service Pack demora muito para ser lançado pela Microsoft, e, normalmente, traz mudanças significativas em vários aspectos para o sistema operacional. Para se ter uma ideia, o Windows 7 atualmente está com o SP1 (lançado em setembro de 2010) e até agora, não houve o SP2 (estamos no início de 2013). O Windows XP, por sua vez, chegou ao SP3!

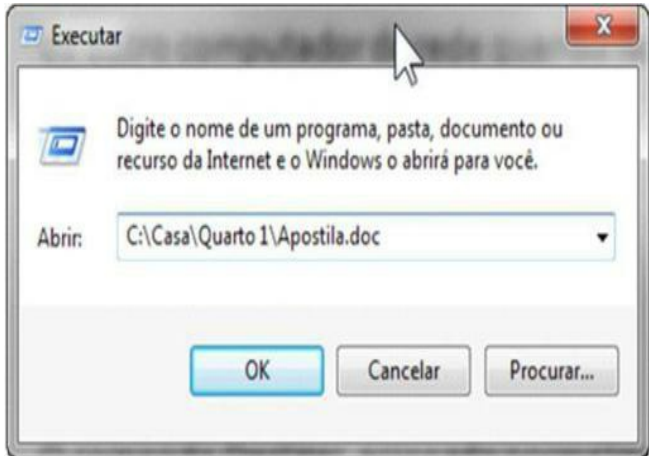
Só um lembrete: algumas dessas atualizações exigem que se reinicie o computador para que tenham efeito (pois algumas delas só se instalam ou no momento do desligamento, ou no momento da inicialização do Windows). Você reconhece que há atualizações que precisam do reinício do computador por meio de um ícone (um escudo amarelo) no comando Desligar do menu Iniciar.



*Figura 4.126 – Atualizações esperando reiniciar para instalar.*

#### 4.5.4. Comando Executar

Permite ao usuário abrir qualquer arquivo (executável ou de dados) e pasta desde que se conheça o endereço completo para achar o referido objeto. O comando Executar é encontrado no menu Iniciar (Iniciar/Todos os Programas/Acessórios/Executar). Verifique, a seguir, a janela do comando Executar em ação.



*Figura 4.127 – Comando Executar, apontando para um arquivo DOC.*

Note a necessidade de escrever o caminho completo para o objeto que se deseja abrir (seja um programa, um arquivo, uma pasta ou uma página da Internet).

É possível, inclusive, executar arquivos em outros computadores da rede, bastando informar \\computador\diretórios\arquivo. O símbolo “\\” precede o nome de outro computador da rede quando fazemos referência a ele (já visto anteriormente).

#### **4.5.5. Comando Desligar**

O comando para desligar o computador é acionado normalmente por meio de uma caixa de listagem na parte inferior direita do menu Iniciar. Além do comando Desligar, em si, é possível localizar, neste menu, outras opções.



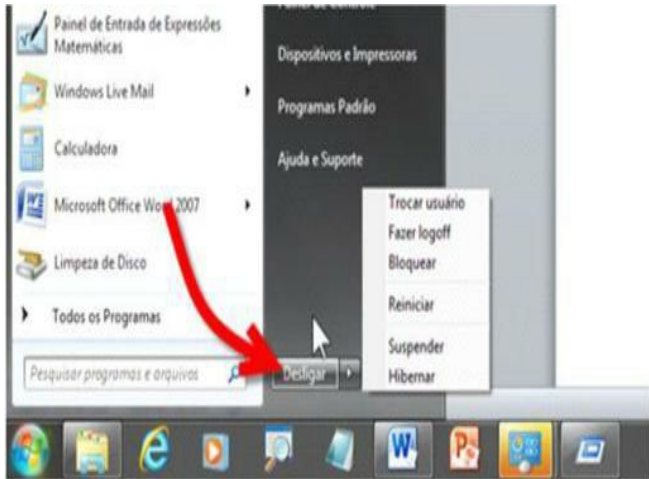


Figura 4.128 – Comando Desligar do Windows 7 e suas opções.

As opções apresentadas nesta caixa são:

- **Desligar:** o computador será desligado;
- **Reiniciar:** o computador será desligado e religado imediatamente;
- **Suspender:** coloca o computador em estado suspenso (estado de baixo consumo de energia: monitor, discos rígidos e outros equipamentos são desligados, mas o sistema continua sendo executado na memória principal). Lembre-se de que neste estado, o **processador** e a **memória principal** continuam funcionando (além, claro, da placa-mãe), mas os circuitos desnecessários são desligados. Ou seja, o micro **continua ligado!**
- **Hibernar:** grava todo o conteúdo da memória principal em um arquivo no disco rígido e, em seguida, **desliga o computador**. Quando o computador for religado, o Windows vai ler o conteúdo desse arquivo e jogá-lo na memória RAM, para que o computador reinicie exatamente do mesmo ponto em que havia parado.

Lembre-se de que a hibernação criará um arquivo, do mesmo tamanho da memória principal física (a RAM), em uma unidade de disco rígido (normalmente a unidade C:). Ao reiniciar o computador, todas as janelas que estavam abertas e todos os textos que estavam sendo vistos no

momento do desligamento do micro serão recuperados exatamente da mesma forma como estavam no momento da hibernação.

- **Fazer Logoff:** solicita o fechamento de todos os programas ativos e desloga (desconecta) o usuário atual (ou seja, quem estiver usando o micro neste momento) sem desligar a máquina. O Windows retornará para a tela de logon (para esperar pelo NOME e SENHA de algum usuário). É como se você “batesse o ponto” esperando pelo funcionário que irá lhe substituir.

- **Bloquear:** aciona a tela de bloqueio do computador. É uma tela semelhante à tela de logon, pedindo senha do usuário atual. Neste modo, os programas não são fechados, mas o computador é bloqueado (protegido contra uso) e fica esperando o desbloqueio (por parte do usuário atual) ou a troca de usuário (abertura de outro usuário simultaneamente), se estiver disponível.

- **Trocar Usuário:** permite que outro usuário faça o logon (habilite-se para trabalhar) no computador sem fechar ou deslogar o usuário atual. Sim! Os dois usuários ficam abertos e ativos em posições diferentes da memória RAM do computador.

É claro que eles não podem USAR o micro simultaneamente (só há um teclado, um mouse e um monitor, né?). Mas os dois usuários ficam abertos ao mesmo tempo, permitindo que, quando um sair, o outro assuma imediatamente, com o acionamento do comando Trocar Usuário.

Quanto um está ativo, em primeiro plano, sendo usado, o outro usuário fica em segundo plano, inativo, por enquanto, mas com todos os programas que deixou abertos lá, intactos!

Acionar a combinação de teclas **CTRL + ALT + DEL** também permite acesso a uma janela que contém as todas as opções da caixa Desligar, além de opções para troca da senha (do usuário atual) e acesso ao Gerenciador de Tarefas.

O Gerenciador de Tarefas é um utilitário, pertencente ao Windows, que permite manusear, entre outras coisas, os programas em execução no computador (programas que estão abertos na memória RAM). É possível, inclusive, excluir um programa da RAM forçadamente (se ele estiver travado, por exemplo, ou seja, “não estiver respondendo”).

Use, para isto, a opção **Finalizar Processo**, no botão desta janela!

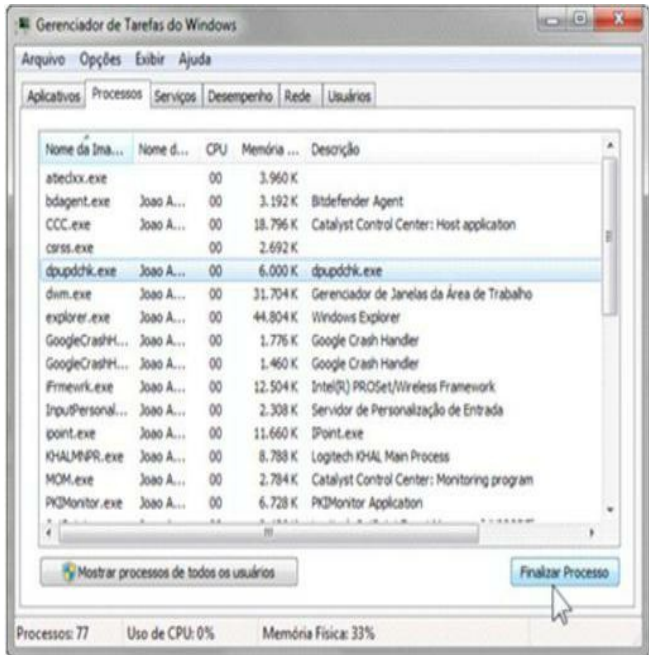


Figura 4.129 – Gerenciador de Tarefas do Windows 7.

#### 4.5.6. Registro do Windows (Registry)

O Windows, assim como todo sistema operacional, é formado por diversos arquivos que guardam suas opções de funcionamento. O conjunto mais importante de informações do Windows é, sem dúvidas, o **Registro**.

O Registry, ou registro, é um banco de dados como todas as informações de configuração do Windows, desde o papel de parede até o perfil de cada usuário do computador. Um usuário não

tem motivos para mexer no registro, em vez disso, os programas, quando instalados ou desinstalados e o próprio Windows, quando tem alguma configuração alterada, fazem alterações no registro.

O programa usado para alterar o registro manualmente é o **Regedit** (Editor do Registro), que pode ser executado através da digitação de seu nome (regedit.exe) no comando Executar ou no campo de pesquisa do menu Iniciar.

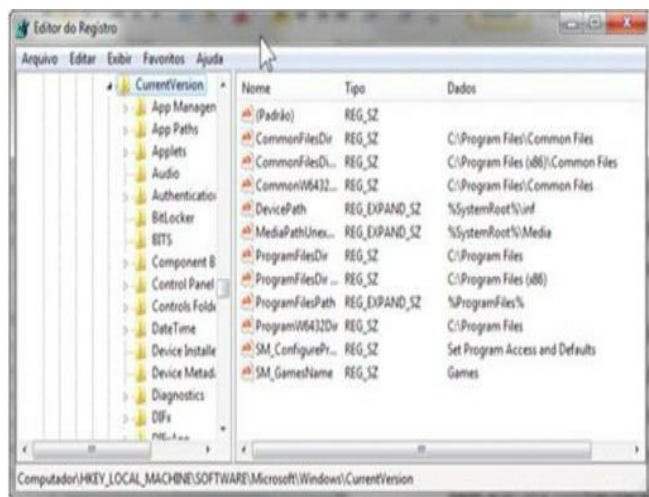


Figura 4.130 – O programa Editor de Registro do Windows (Regedit).

Raramente temos de nos preocupar em alterar configurações no registro, e, sinceramente, é bom que não seja necessário mesmo! Qualquer alteração errada pode fazer o sistema operacional parar de funcionar completamente.

Um exemplo de quando um usuário deve mexer diretamente no registro é quando alguns malwares (programas maliciosos, como vírus, cavalos de Troia, spywares etc.) infectam o sistema. Uma das primeiras coisas que um malware faz é garantir que seja executado sempre que o computador for ligado e, para isso, a maioria dos malwares altera o registro para informar ao Windows que, quando o sistema for ligado, o malware também seja executado.

Para corrigir isso, um usuário pode alterar manualmente o registro na posição correta para que a configuração do malware deixe de funcionar (ou seja, o usuário vai ter de apagar o que o malware escreveu no registro). Não preciso dizer que o usuário precisará saber exatamente o que foi que o malware fez, não é?

#### 4.5.7. A Estrutura de pastas do Windows 7

Quando o Windows 7 é instalado, seus arquivos são copiados para uma das unidades de disco rígido do computador (normalmente a primeira delas, que será chamada, pelo Windows, de C:) e criará, nessa unidade, algumas pastas para guardar esses arquivos.

Quando instalado, o Windows já cria, automaticamente, na unidade C, as pastas Program Files, Program Files (x86), Users e Windows, entre outras. Vamos conhecer um pouco mais acerca desses diretórios (pastas).

##### 4.5.7.1. Program Files e Program Files (x86)

A pasta Program Files serve para guardar pastas correspondentes aos vários programas instalados no seu computador. Sim! Os aplicativos (programas), quer tenham vindo com o Windows, quer tenham sido instalados posteriormente, ficam, normalmente, nesta pasta.

Cuidado: a pasta Program Files, em uma edição do Windows de 64 bits (como é o meu caso), serve para armazenar apenas os aplicativos também feitos em 64 bits.

Os aplicativos (e demais programas) construídos para 32 bits (o sistema operacional Windows de 64 bits consegue executar aplicativos de 32!) ficam armazenados na pasta ***Program Files (x86)***.



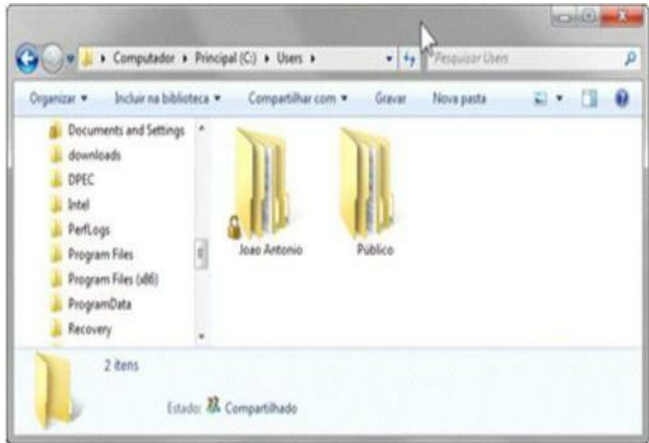
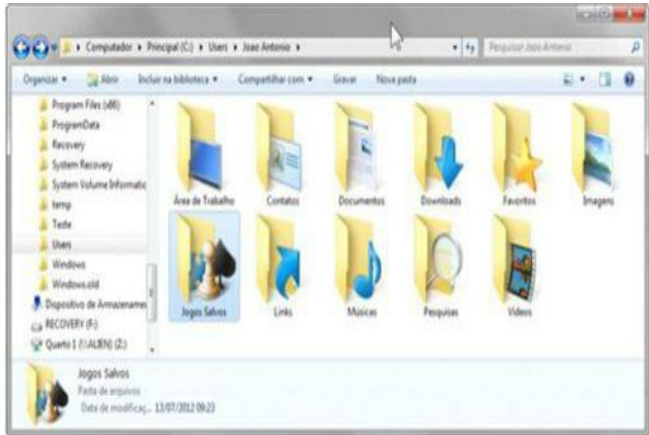


Figura 4.132 – A pasta “Users” e as pastas dos usuários do sistema (só tem um usuário).

A pasta **João Antonio** pertence a um usuário com esse nome (dããã! Sou eu, né?) e só poderá ser acessada por este usuário. Já a pasta **Público** é criada para colocar componentes (arquivos e pastas) que serão acessados por todos os usuários daquele computador.

A pasta João Antonio é chamada “**pasta pessoal**” do usuário João Antonio.

Vamos dar uma olhada nas pastas que o Windows 7 cria dentro da pasta pessoal do usuário? Perceba como tudo é “organizadinho”, “bonitinho”, “arrumadinho”.



*Figura 4.133 – Pastas dentro da pasta pessoal do usuário João Antonio.*

Só um lembrete: você não precisa “obedecer” à organização que o Windows deu normalmente! Essas pastas são apenas uma ideia, ou seja, uma sugestão. Claro que elas estão todas tão “arrumadinhas” que a gente se sente até “mal” de não aceitar tamanha deferência, né?

Você pode criar suas próprias pastas, onde quer que seja! Com o nome que quiser! Inclusive, pode criar aí mesmo nessa pasta pessoal! Afinal, ela é sua!

#### **4.5.7.3. Windows**

Esta pasta guarda os arquivos de configuração e de programas principais do próprio sistema operacional Windows 7. Quando a “bronca” acontecer no Windows, o arquivo danificado provavelmente estará aqui.



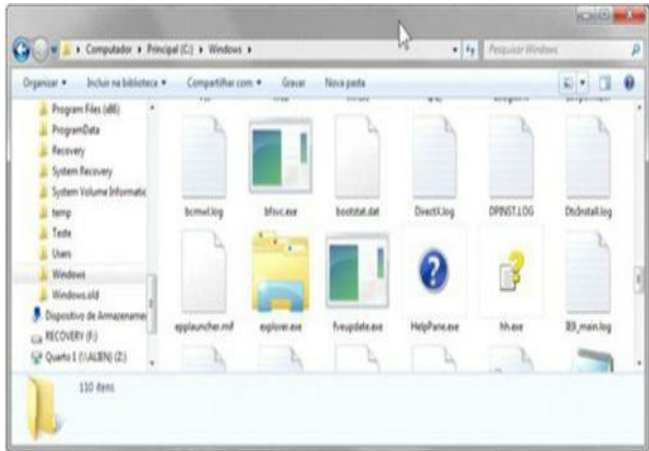


Figura 4.134 – Pasta Windows (olha lá o Windows Explorer!).

Há muito conteúdo importante e “delicado” na pasta Windows. Mexer nos arquivos aqui dentro é algo pouco recomendado para os usuários com menos experiência (e até mesmo para alguns que acham que têm experiência).

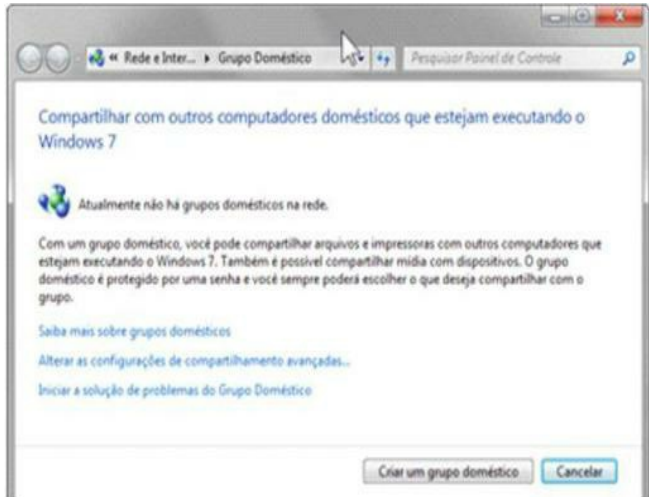
#### 4.5.8. Grupo Doméstico

Um dos principais novos recursos do Windows 7 é, sem dúvidas, o Grupo Doméstico. Essa novidade traz inúmeras facilidades para quem quer compartilhar recursos em rede (em casa).

Para se criar um grupo doméstico, é necessário que:

- a. haja uma rede de computadores reconhecida como **rede doméstica** (explico isso adiante).
- b. todos os computadores utilizem Windows 7.

Para criar o grupo doméstico, basta acionar a opção Grupo Doméstico dentro do Pannel de Controle (categoria Redes e Internet), ou digitar “Grupo Doméstico” no campo de pesquisa do menu Iniciar. A janela a seguir é aberta:



*Figura 4.135 – Criando um Grupo Doméstico.*

Clicando no botão Criar um grupo doméstico, você será enviado para a janela que pergunta que tipo de documentos você deseja compartilhar (na verdade, quais bibliotecas do seu computador), ou seja, quais bibliotecas serão vistas pelos outros computadores do seu grupo doméstico.



*Figura 4.136 – Escolhendo as bibliotecas (e impressoras) compartilháveis.*

Depois disso, uma senha (código único) é mostrada a você pelo Windows. Esse código deverá ser digitado em cada computador que deseje ingressar neste grupo doméstico.

## Use esta senha para adicionar outros computadores ao grupo doméstico

Antes de acessar impressoras e arquivos localizados em outros computadores, adicione esses computadores ao seu grupo doméstico. Você precisará da senha a seguir.

Anote sua senha:

aA5ds7woaA

[Imprimir senha e instruções.](#)

Se esquecer a senha do grupo doméstico, você poderá visualizá-la ou alterá-la abrindo Grupo Doméstico no Painel de Controle.

[Como outros computadores podem ingressar no meu grupo doméstico?](#)

Concluir

Figura 4.137 – Senha definida para o Grupo Doméstico.

Em resumo: um grupo doméstico é criado em um dos computadores da rede e, em seguida, todos os demais ingressam (passam a fazer parte) do grupo por meio da senha gerada no computador que o criou.

Nos demais computadores da rede, basta acionar Grupo Doméstico (no Painel de Controle, também) e a janela apresentada será outra: a de Ingressar no Grupo Doméstico (o Windows 7 dos demais micros irá “detectar” a presença de um grupo já criado).



*Figura 4.138 – Ingressar no grupo doméstico.*

Depois, é só digitar a senha ofertada pelo primeiro micro e este computador passará a fazer parte do grupo doméstico recém-criado.



*Figura 4.139 – Digite a senha do Grupo.*

Finalmente, para ter acesso ao grupo doméstico, basta, no Painel de Navegação do Windows Explorer, acessar o item Grupo Doméstico. É muito fácil! Lá estará listado o usuário que aceitou entrar no grupo doméstico e suas bibliotecas (aquelas que ele escolheu compartilhar).



Figura 4.140 – Visualizando dois Micros no Grupo Doméstico.

No exemplo acima, temos dois computadores conectados no grupo doméstico (além, é claro, do meu que estou usando agora): TOUCHHP, com o usuário Be-a-Byte e MOBALIEN, tendo o usuário Joao Antonio logado no momento.

#### 4.5.8.1. O que é uma Rede Doméstica?

Quando um computador se conecta a uma rede pela primeira vez, ele é questionado acerca do “tipo” daquela rede (o “local da rede”, que determina a relação que aquele computador deve ter com aquela rede e com os demais micros dela).



### Rede doméstica

Se todos os computadores da rede estiverem na sua casa e puderem ser reconhecidos, você estará em uma rede doméstica confiável. Não escolha esta opção para lugares públicos, como restaurantes ou aeroportos.



### Rede corporativa

Se todos os computadores da rede estiverem em seu local de trabalho e puderem ser reconhecidos, você estará em uma rede corporativa confiável. Não escolha esta opção para lugares públicos, como restaurantes ou aeroportos.



### Rede pública

Se não reconhecer todos os computadores da rede (por exemplo, quando estiver em um restaurante ou no aeroporto, ou quando estiver usando banda larga móvel), você estará em uma rede pública, não confiável.

Figura 4.141 – Definindo o Local da Rede (quando se conecta pela primeira vez!).

São três as opções de “Local da Rede”:

- **Rede Doméstica:** deve ser escolhida se você está em sua casa, com micros confiáveis e conhecidos. Esta opção permite muito mais “liberdade” no compartilhamento de recursos. É a única que permite a criação do Grupo Doméstico.
- **Rede Corporativa:** use essa opção caso seu computador esteja se conectando a uma rede na empresa. Haverá possibilidade de compartilhamento de recursos, sim, mas bem mais



restritos que na rede doméstica.

- **Rede Pública:** use essa opção quando estiver conectado a uma rede desconhecida, como de hotéis, aeroportos, shoppings e restaurantes. Essa opção permite o acesso a Internet por meio da rede em questão, mas limita e proíbe (o que é o certo) a maioria dos recursos de compartilhamento.

Para conferir em que local de rede você está conectado (e até mesmo alterar esse local), pode-se acessar o item **Central de Rede e Compartilhamento**, existente na categoria **Rede e Internet** do Painel de Controle. Lembre-se de que você pode acessar qualquer item do Painel de Controle por meio do campo de pesquisa do menu Iniciar (para facilitar):

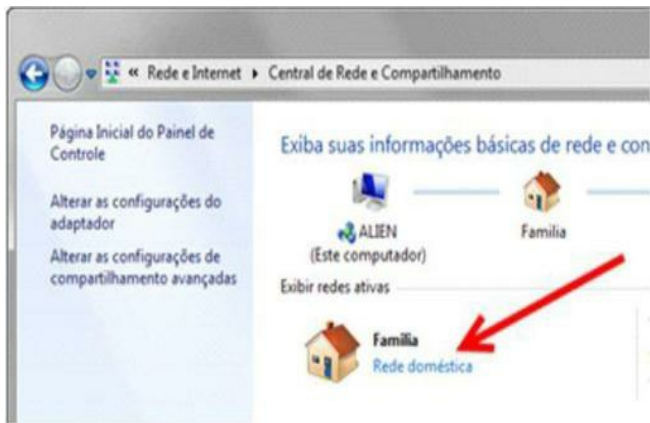


Figura 4.142 – Visualizando (e podendo redefinir) o Local da Rede.

#### 4.6. Windows 8 – O mais novo!

A Microsoft já lançou, em 2012, a versão mais atual de seu sistema operacional: o Windows 8. Com uma interface totalmente renovada, premiando os tablets e outros dispositivos com touchscreen (tela sensível ao toque), este Windows, muito provavelmente, não será palco de questões em prova durante muito tempo!

Portanto, não há, pelo menos por agora, por que se preocupar com este assunto! Segue, porém, uma pequena amostra da tela inicial deste programa!



*Figura 4.143 – Windows 8 – ideal para micros que funcionam como tablets.*

É sério, caro leitor! Esse programa não é tão “interessante” para concursos públicos porque ele é mais voltado para os dispositivos que usam touchscreen, coisa que não se encontra com tamanha facilidade em serviço público (logo, não há motivo para cobrá-lo em prova!).

#### 4.7. Questões de Windows

1. No MS-Windows 7, a operação de exclusão definitiva de um arquivo, sem movê-lo para a lixeira, deve ser acompanhada do pressionamento da tecla:
  - a) Scroll;
  - b) Ctrl;
  - c) Alt;
  - d) Shift;
  - e) Tab.
2. No Windows 7, em sua configuração padrão e original:
  - a) a ativação do Firewall do Windows é feita por intermédio do menu Arquivo do Windows Explorer;
  - b) a opção de desligamento automático do monitor, após um determinado tempo, está disponível em Opções de energia no Pannel de controle;
  - c) não é possível a renomeação de um nome de arquivo clicando sobre o nome do arquivo com o botão direito do mouse;
  - d) após sua instalação, a alteração de uma conta de usuário fica bloqueada permanentemente para uso;
  - e) não é possível a exclusão de um arquivo clicando sobre o nome do arquivo com o botão direito do mouse.
3. Para controlar o consumo de energia, o Windows pode colocar o computador, após um determinado período de inatividade, em modo de:
  - a) hibernação, que mantém o conteúdo da RAM, desliga a maioria dos circuitos e não permite desconectá-lo da rede elétrica;
  - b) suspender, que mantém o conteúdo da RAM, desliga o computador e não permite desconectá-lo da rede elétrica;
  - c) suspender, que transfere o conteúdo da RAM para o HD, desliga a maioria dos circuitos e não permite desconectá-lo da rede elétrica;
  - d) suspender, que transfere o conteúdo da RAM para o HD, desliga o computador e permite desconectá-lo da rede elétrica;
  - e) hibernação, que transfere o conteúdo da RAM para o HD, desliga o computador e permite desconectá-lo da rede elétrica.
4. Considerando que o sistema operacional Windows apresenta configurações padrão de arquivos, temos que a extensão:
  - a) “.xls” refere-se a um arquivo do Microsoft Excel;
  - b) “.doc” refere-se a um arquivo do Microsoft Access;
  - c) “.zip” refere-se a um arquivo padrão texto;
  - d) “.bmp” refere-se a um arquivo de música;
  - e) “.exe” refere-se a um arquivo de imagem.

5. Para minimizar todas as janelas abertas atualmente no Windows 7, pode-se acionar um clique no botão:
- a) Minimizar Todos, ao lado direito da área de Notificação;
  - b) Área de Trabalho, no menu Iniciar;
  - c) Área de Trabalho, na Área de Notificação;
  - d) Mostrar Área de Trabalho, no menu Iniciar;
  - e) Mostrar Área de Trabalho, ao lado direito da área de Notificação.
6. Acerca do Sistema Operacional Microsoft Windows e de seus aplicativos, julgue os itens a seguir.
- I. O programa Paint, do Windows 7, é utilizado para editar imagens de bitmap. O formato padrão de arquivo salvo por esse programa é o PNG, mas os formatos JPG e GIF, muito usados na Internet, também são suportados.
- II. Os arquivos apagados de um disco rígido ou de um CD-ROM são enviados para a lixeira, de onde podem ser recuperados pelo usuário através de procedimentos simples, mas os arquivos do disquete não são enviados para a lixeira.
- III. O campo Pesquisar, do menu Iniciar, permite que se encontrem itens do Painel de Controle usando como critério parte do nome dos itens.
- IV. CTRL+TAB permite acionar o recurso de Flip 3D no Windows 7.
- Os itens que apresentam todas as assertivas corretas são:
- a) I e II;
  - b) II e III;
  - c) III e IV;
  - d) I e III;
  - e) II e IV.
7. Acerca do programa desfragmentador de disco, presente no sistema Windows, é correto afirmar que:
- a) pode ser usado para localizar arquivos e pastas que estão em locais diferentes, espalhados pelo disco rígido;
  - b) desfragmenta CDs e disquetes também, além do HD;
  - c) pode ser usado em discos com FAT32 e NTFS;
  - d) não pode ser usado em discos formatados com o sistema FAT32;
  - e) apaga o conteúdo da FAT quando é executado.

## 5.1. Conceito de aplicativos

Aplicativos são os programas de computador criados para solucionar problemas dos usuários da informática. Um processador de texto, uma planilha eletrônica, um programa para construir mapa astral, são exemplos de aplicativos.

Há vários aplicativos cobrados em provas, alguns deles, claro, vão ser mais aprofundados ao longo deste livro (em capítulos posteriores)! Mas vamos dar uma visão geral (e superficial) acerca da maioria deles neste capítulo.

### 5.1.1. Tipos de aplicativos

Há algumas classificações possíveis no universo de aplicativos, vamos a algumas delas:

- **Processador de Texto:** programa com a função de permitir que o usuário construa os mais trabalhados documentos de texto profissionais, desde cartas e bilhetes, passando por relatórios, apostilas e livros. O Microsoft *Word* e o LibreOffice *Writer* são exemplos desse tipo de programa.
- **Planilha Eletrônica:** software que auxilia o usuário na tarefa de criar e manipular dados numéricos em tabelas. Normalmente esses programas também fornecem recursos para a construção de gráficos a partir dessas tabelas numéricas. O *Excel* (da Microsoft) e o *Calc* (do conjunto LibreOffice) são representantes dessa classificação.
- **Gerenciador de Bancos de Dados:** é o programa que manipula dados em estruturas organizadas chamadas bancos de dados. Normalmente utilizado em sistemas de controle de estoque e cadastro de clientes das empresas. O Microsoft *Access* é um exemplo e o LibreOffice *Base* é seu principal concorrente.
- **Gerenciadores/Editores de Apresentações de Slides:** são programas que permitem a construção de apresentações de slides, normalmente usadas em palestras e aulas. O Microsoft *PowerPoint* é o mais famoso deles! O conjunto de programas LibreOffice também tem o seu: o *Impress*.

Existem muitas outras classificações que, por não serem unanimidade entre os autores, não serão vistas aqui.

## 5.2. Instalação de um programa

Quando um determinado programa não pertence ao sistema operacional, ele deve ser adicionado ao computador através de um processo chamado **Instalação**.

A instalação consiste em um processo de cópia dos arquivos que formam o programa em questão (ou parte dele) para o disco rígido do computador e, além disso, um registro no sistema operacional sobre a existência do novo software (alteração no registry do Windows).

Funciona mais ou menos assim: quando um usuário quer instalar um novo jogo, por exemplo, ele deve inserir a unidade de disco em que estão os arquivos do jogo (normalmente um CD,

DVD ou pen drive) e iniciar o processo de instalação (que, quase sempre, é executado por um programa instalador).

Depois de completo o processo de instalação, o jogo estará completamente (ou quase) copiado para o HD da máquina em que foi instalado, e o sistema operacional reconhece que o programa existe, então ele pode ser utilizado sempre que o usuário o execute (duplo clique no seu ícone).

Note: além de poder vir em DVD ou outra mídia, é comum também o ato de **baixar** (copiar da internet) o **arquivo instalador** (um EXE, normalmente). Esse único arquivo, quando executado, irá proceder com o processo de instalação do programa a que se refere.

### 5.2.1. Desinstalação de um programa

Desinstalar um programa é um processo tão fácil quanto instalar, requer apenas que o usuário localize o programa **desinstalador** (que normalmente acompanha o aplicativo) e acione-o, deixando tudo a cargo do próprio programa desinstalador.

Outra forma muito segura é usar o ícone **Programas e Recursos**, do Painel de Controle do Windows. Esse componente apresenta uma listagem de todos os programas instalados no computador e registrados no sistema operacional Windows.

Basta escolher o programa desejado e acionar o comando da remoção.

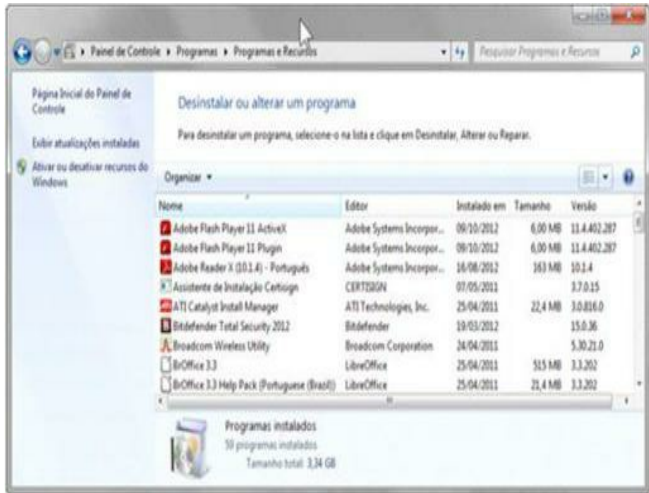


Figura 5.1 – Janela do item Programas e Recursos, do Painel de Controle.

Atenção: durante a desinstalação de um programa, seus arquivos, que estão gravados no disco rígido, são removidos dessa unidade. Não é necessário, portanto, excluir os arquivos manualmente após o processo de desinstalação.

Outra informação importante: não se deve simplesmente apagar os arquivos de um programa (lá na pasta dele) manualmente, julgando que isso constitui o processo de desinstalação. Desinstalar um programa é dizer ao Windows que o programa não está mais em funcionamento nesse computador; portanto, deve-se seguir o procedimento correto.

### 5.3. Classificação quanto à Licença de Uso

Um programa de computador pode ser classificado de algumas maneiras, no que se refere ao direito de uso que um determinado usuário tem sobre ele.

- **Software:** é, no geral, uma classificação que envolve todo tipo de programa, mas pode significar também, para alguns autores, os programas pagos, programas pelos quais se deve pagar uma taxa chamada Licença de Uso.

Outros termos que deixam claro que a classificação é a de software pago são: “Software

Proprietário” ou “Software Comercial”.

- **Freeware:** são os programas completos e gratuitos, que os programadores criam e distribuem, sem custo, para os usuários. Por que fazem isso? Não me pergunte! Por serem revolucionários, reacionários, subversivos, underground, do PSTU, sei lá! ;-)
- **Shareware:** são programas distribuídos gratuitamente pelos seus fabricantes, mas que não são completos. Normalmente esses programas possuem pequenas limitações de uso (tempo, menos recursos etc.), que apenas dão o “gostinho” ao usuário, que, se quiser mais, tem de comprar a versão comercial.

Há shareware e freeware de diversos tipos disponíveis na Internet, desde jogos até mesmo grandes sistemas de controle de clientes e estoque.

Os dois programas estudados a seguir não acompanham o sistema Windows; portanto, devem ser adquiridos de forma separada (normalmente por downloads da Internet). São eles:

- WinZip;
- Adobe Acrobat Reader.

#### 5.4. WinZip – Compactador de arquivos

O WinZip é um programa shareware com a função de compactar arquivos. Seu funcionamento é simples: um ou vários arquivos de qualquer tipo são “prensados” e se tornarão um único arquivo (com extensão ZIP) que ocupará uma menor quantidade de bytes na unidade de disco.

Com o WinZip, um ou mais arquivos são compactados e colocados em um único arquivo resultante.

“Ei, João, qual a razão de se compactar um ou mais arquivos?”

Enviar por e-mail arquivos menores seria uma boa razão. Embora, hoje em dia, a Internet seja demasiadamente mais rápida que há alguns anos, ainda assim, compactar arquivos é uma “boa pedida” para transportá-los por meio do correio eletrônico.

“Como uso um arquivo compactado?”

Com o próprio WinZip é possível descompactar um determinado arquivo ZIP, extraindo seu conteúdo para poder ser usado novamente.

Repito: normalmente a compactação é utilizada para facilitar o transporte da informação por meios que exigem tamanhos menores. Para que o arquivo compactado seja usado, deve-se extrair de seu interior os dados que nele foram armazenados.

Quando compactamos arquivos, não podemos esperar que todos eles reduzam de tamanho de forma igual ou proporcional. Cada tipo de arquivo apresenta uma taxa de redução diferente, alguns arquivos são mais reduzidos (já encontrei taxas de até 95% de redução), outros quase não se reduzem. Textos e planilhas normalmente são mais reduzidos, as fotos e desenhos tendem a não diminuir muito seu tamanho em bytes.

Há diversos outros programas para compactar e descompactar arquivos como o WinRAR e o WinAce. O próprio Windows 7 consegue entender os arquivos com extensão ZIP como se fossem pastas (mostra-se “pasta compactada” no ícone do arquivo ZIP). Portanto, para o Windows 7, não é necessário possuir o WinZip (ou qualquer outro programa) para poder extrair ou compactar arquivos zipados.



## 5.5. Adobe Reader

É um programa freeware, fabricado pela empresa Adobe (famosa por programas de desenho e fotografia, como o Adobe Photoshop, considerado o melhor programa de edição de fotografia da atualidade). O Acrobat Reader tem a função de ler (visualizar) arquivos no formato PDF. A tecnologia dos arquivos PDF foi desenvolvida pela própria Adobe; portanto, o formato PDF é uma propriedade dela.

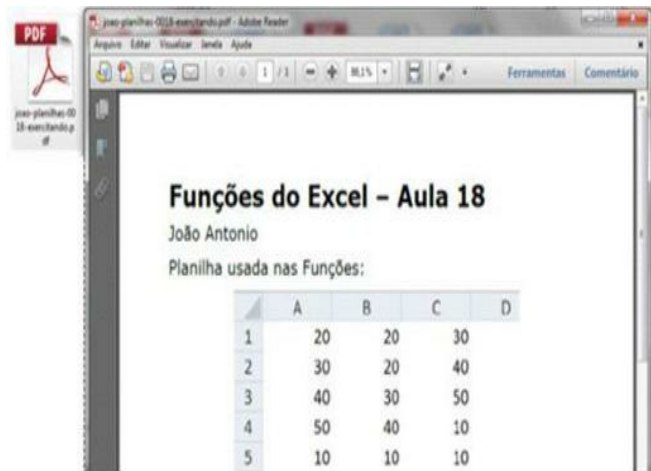


Figura 5.2 – Ícone do arquivo PDF e uma janela do programa Adobe Acrobat Reader em ação.

Para que serve um arquivo PDF? É um documento gráfico (ou seja, permite textos e figuras) que não pode ser (tão facilmente) alterado pelo usuário. Atualmente é muito comum, na Internet, uma empresa disponibilizar seus materiais didáticos, técnicos e de referência nesse formato de arquivo, com a grande vantagem de ter a certeza (ou quase) de que o arquivo não será editado por pessoas não autorizadas.

São vantagens dos arquivos PDF:

- Não podem ser modificados (facilmente);
- São arquivos menores (em bytes) que os arquivos originais (embora alguns arquivos DOCX consigam ficar menores que os respectivos PDFs).

c. São vistos e impressos por qualquer computador da mesma forma como foram criados (páginas, cores, fontes, imagens), independentemente do programa que foi usado para criar o arquivo que deu origem ao PDF. Basta o usuário possuir o Acrobat Reader para poder ler e, se quiser, imprimir o PDF (se, claro, não estiver bloqueado para impressão).

“Ei, João, como crio um arquivo PDF com o Adobe Reader?”

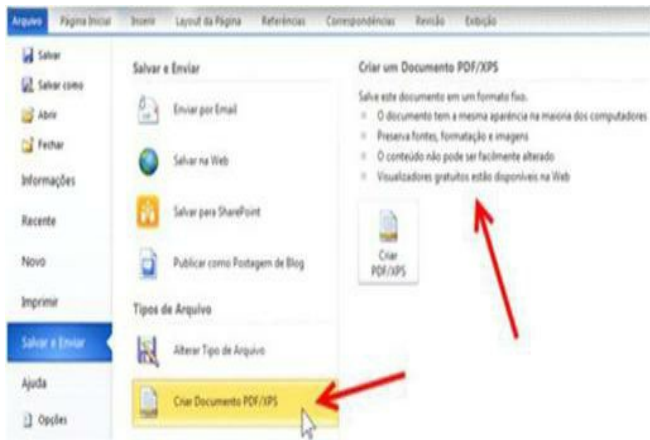
Caro leitor, o Adobe Reader é apenas o programa leitor (não permite criar PDF, só ler) e é **gratuito**. O programa Adobe Acrobat (ou Acrobat Professional) é o programa que pode criar arquivos no formato PDF, mas **não é gratuito** (aliás, é caro pra cacete!).

Se o usuário possui o programa criador (Adobe Acrobat), basta escrever qualquer documento em qualquer programa (Word, Excel, PowerPoint, CorelDRAW etc.) e solicitar que o Acrobat o converta em PDF.

Caso tenha detectado algum erro no arquivo PDF, o usuário tem de descartar o referido arquivo PDF, corrigir o problema no arquivo original (seja ele em que programa for) e solicitar que se converta novamente, criando o arquivo PDF mais uma vez.

Hoje em dia, porém, muitos programas comerciais (como o próprio Word e os demais programas do Microsoft Office – versão 2010) conseguem salvar um arquivo diretamente no formato PDF, tornando desnecessária a aquisição do Adobe Acrobat (o caro pra cacete!).

Note: o Word não lê PDFs (isso é trabalho do Adobe Reader), mas pode criar PDFs (no comando Salvar Como).



Todos os programas do conjunto LibreOffice (BrOffice) também salvam normalmente em PDF (aliás, o faziam há muito mais tempo que o Microsoft Office).

Ahhh... tá cansado de ler “Office” o tempo todo e fica se perguntando o que é isso? Dá uma olhada a seguir... ;-D

## 5.6. Suítes de programas

Chama-se suíte de programas, ou pacote de programas, um conjunto de softwares comercializados juntos, em uma mesma embalagem. Os principais produtos desta categoria são desenvolvidos para trabalho de escritório (programas para uso geral).

### 5.6.1. Microsoft Office

A suíte mais conhecida para nós é o Microsoft Office, que reúne os principais programas para automação de escritório desenvolvidos pela Microsoft, a mesma empresa que desenvolve o Windows. Fazem parte do Microsoft Office:

- **Microsoft Word:** processador de textos;
- **Microsoft Excel:** planilha eletrônica;
- **Microsoft PowerPoint:** Programa para criação e edição de apresentações multimídia (usadas em palestras, por exemplo);
- **Microsoft Access:** gerenciador de bancos de dados (este programa não está presente em todas as versões do Office!);
- **Microsoft Outlook:** central de comunicação que permite o envio e recebimento de e-mails, fax, agenda de reunião etc. (Este programa não está presente em todas as versões do Microsoft Office!);

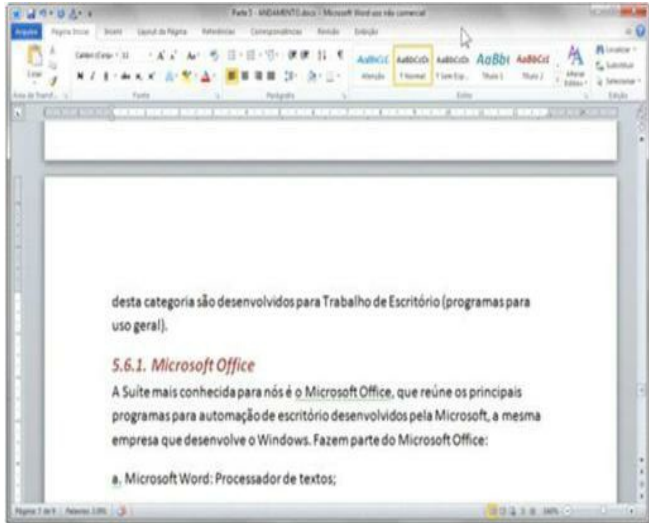


Figura 5.4 – Word 2010 – Editando... Este livro! ;-D

O Microsoft Office traz, ainda, outros pequenos aplicativos para auxiliar o desempenho dos aplicativos principais, mostrados anteriormente. Como exemplo, podemos citar o **Microsoft Equation** (editor de equações), que permite construir equações complexas usadas na matemática.

Além de vários aplicativos, grandes e pequenos, o Office apresenta uma linguagem de programação própria, para tornar os aplicativos mais “personalizáveis” e criar verdadeiros programas com os arquivos do Word, Excel e Access.

A linguagem de programação que acompanha o Microsoft Office chama-se VBA (Visual Basic para Aplicações) e é uma versão reduzida do Visual Basic, linguagem de programação profissional que a Microsoft desenvolve e comercializa.

Posteriormente, vamos conhecer mais detalhes sobre alguns dos principais aplicativos do Office, pois é muito comum encontrá-los em concursos públicos.

Atualmente, o Microsoft Office encontra-se na **versão 2010**, que, por algum tempo, ainda, será a versão cobrada em provas de concursos.

### 5.6.2. LibreOffice (antigo BrOffice)

Alguns programadores ao redor do mundo se juntaram para desenvolver um pacote de programas de escritório livre (acessível e modificável por todos). Dessa iniciativa surgiu o **OpenOffice** (ainda existente – [www.openoffice.org](http://www.openoffice.org)).

Um grupo (dissidente) de programadores brasileiros pegou o OpenOffice e o adaptou à nossas características (nossas, digo, do Brasil) e criou o bastante famoso **BrOffice** (acessível em [www.broffice.org](http://www.broffice.org) – anunciaram lá, oficialmente, o fim do projeto BrOffice).

Em 2011, os programadores do BrOffice passaram a fazer parte de uma equipe mundial que trabalha com o mesmo objetivo: fazer um versão melhorada do OpenOffice (que é mantido por outra equipe) – eis que surgiu, deste esforço, o **LibreOffice**, o pacote livre de programas de escritório que pode vir a ser a “menina dos olhos” das bancas examinadoras nestes próximos concursos (acesse em [www.libreoffice.org](http://www.libreoffice.org)).

(Na verdade, ao que parece, houve uma série de discussões e desentendimentos que fizeram o nome BrOffice não poder mais ser usado, além de, claro, não poder mais ser atualizado, mas isso não importa!)

Ei! Presta atenção! Só para deixar claro... OpenOffice, BrOffice e LibreOffice **não são a mesma coisa!** Mas, por serem “derivados” uns dos outros, eles são muito, mas muitíssimo, parecidos! Logo, comandos, recursos, efeitos são quase iguais (na forma de fazer) nos três, ok?

Hoje em dia, ou você usa o OpenOffice (mantido pela Instituição Apache, dona de outros programas livres conhecidos), ou usa o LibreOffice (esforço de vários programadores ao redor do mundo). Já vi ambos caírem em prova!

“João, posso instalar os dois conjuntos em meu computador?”

Sim, caro leitor! Pode sim! Mas não acredito ter “motivo” para isso! Eles são praticamente idênticos. O que se aprende em um deles, se aprende no outro, basicamente!

Eu acho mais “objetivo” e recomendado instalar o LibreOffice, com o pacote de idioma Português do Brasil, que é o que mais se aproxima do que foi o BrOffice. Aliás, veja uma tela da última versão chamada BrOffice (a versão 3.3):



Texto



Desenho



Planilha



Banco de dados



Apresentação



Fórmula



Abrir...



Modelos...



Figura 5.5 – BrOffice 3.3 – a última versão com esse nome!

E, usando o LibreOffice (atualmente na versão 3.6) como parâmetro, vamos conhecer os principais programas desta suite de escritório:

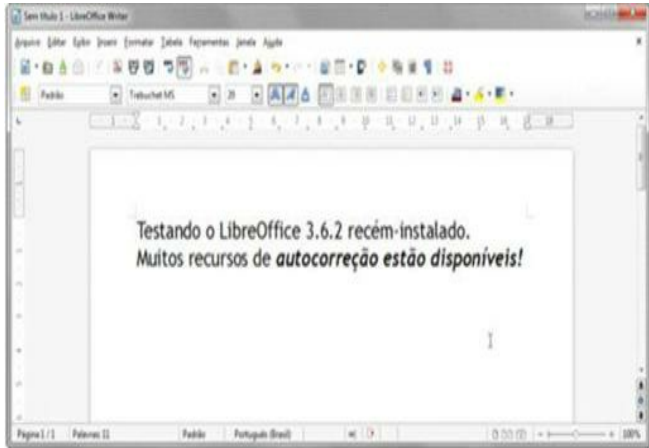
- **LibreOffice Writer:** processador de textos (concorrente do Word);
- **LibreOffice Calc:** planilhas de cálculos (concorrente do Excel);
- **LibreOffice Impress:** apresentações de slides (sonha ser concorrente do PowerPoint);
- **LibreOffice Base:** gerenciador de bancos de dados (para concorrer com o Access);
- **LibreOffice Draw:** programa para desenho vetorial (não há equivalente no Microsoft Office);
- **LibreOffice Math:** programa para desenho de equações matemáticas (seria equivalente ao pequeno “Equation” do Microsoft Office).

Para baixar a última versão do LibreOffice, visite o site da instituição ([www.libreoffice.org](http://www.libreoffice.org)) e clique no botão Download LibreOffice. Se tudo der certo, ele vai reconhecer que você está

usando um navegador com linguagem Português (do Brasil) e sugerirá esta linguagem para o produto a ser baixado!



Figura 5.6 – Home page do projeto LibreOffice.



*Figura 5.7 – LibreOffice Writer 3.6.*

Infelizmente, neste livro, não abordaremos o LibreOffice, até pelo tamanho que o livro teria com esse conteúdo todo! Veremos os dois principais programas do Microsoft Office (o Word 2010 e o Excel 2010).

Conteúdos sobre o BrOffice podem ser adquiridos, posteriormente, na forma de materiais disponíveis no hotsite deste livro na Campus/Elsevier ([www.elsevier.com.br](http://www.elsevier.com.br)) e/ou na seção Materiais Avulsos do Eu Vou Passar ([www.euvoupassar.com.br](http://www.euvoupassar.com.br)).

Também não veremos o PowerPoint (mesmo pertencendo ao Microsoft Office) neste livro. Procure por material deste programa nos endereços citados!



## 6.1. Conhecendo o Microsoft Word

O Word é um programa processador de textos desenvolvido e comercializado pela Microsoft. Esse programa é distribuído dentro do pacote Microsoft Office.

As versões mais recentes do Word são:

- **Word 95 (ou Word 7):** primeira versão deste programa para o ambiente Windows 95.
- **Word 97 (ou Word 8):** a partir deste ponto, os arquivos do Word 97 receberam uma mudança crucial em seu formato interno.
- **Word 2000 (ou Word 9):** não há muitas mudanças internas desde a última versão, apenas detalhes e novos recursos.
- **Word XP (ou Word 10 ou Word 2002):** mais interação com a Internet, mas nada que “assuste” em relação à versão 2000.
- **Word 2003 (ou Word 11):** apareceram mais alguns recursos de formatação, a “interface” está mais parecida com o Windows XP, ou seja, mais “frescuras”.
- **Word 2007 (ou Word 12):** assim como o restante dos programas do Office, o Word 2007 trouxe uma enorme diferença em relação às versões anteriores. A mudança é realmente radical.
- **Word 2010 (ou Word 14 – sim! 14... Pularam o 13! Superstição, será?):** esta é a versão mais recente do Microsoft Word, e é a versão que iremos analisar neste livro (com algumas diferenças do Word 2007, que também pode aparecer em provas!).

“Certo, João: há várias versões do Word... Mas qual devo estudar para a prova? Qual a mais exigida?”

Essa é uma pergunta um tanto difícil, caro leitor. O Word 2003 ainda pode ser cobrado em provas, mas não é o mais comum. É possível, mas menos provável, que se cobre o Word 2007... Mas... Aposto que quase todos os concursos vindouros cobrem Word 2010 (que está sendo o mais usado atualmente, pois veio para substituir, já desde 2010, as versões anteriores).

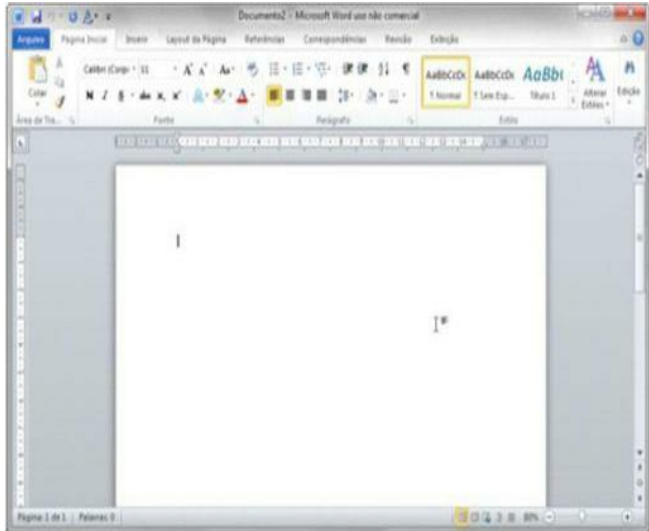
Com relação ao Word 2007, ele é bem semelhante ao Word 2010, e quem aprende um deles basicamente aprende o outro! Vamos abordar as diferenças, quando elas aparecerem, ok?

Sobre o Word 2003, uma apostila gratuita está sendo disponibilizada no site da Editora Campus/Elsevier, na página específica sobre este livro. Esta apostila também pode ser conseguida no Eu Vou Passar ([www.euvoupassar.com.br](http://www.euvoupassar.com.br)).

Versões anteriores do Word podem ser exigidas... Mas é pouco provável!

### 6.1.1. Interface do Word

O Word apresenta sua área de trabalho como uma página em branco pronta para ser preenchida:



*Figura 6.1 – Janela do Word 2010.*

Além dos componentes comuns a qualquer janela (como barra de título, botões minimizar, maximizar e fechar, bordas etc.), há outros componentes próprios da janela do Word 2010 que não podemos deixar de mencionar.

#### **6.1.1.1. Faixa de Opções**

Esse é o nome dado à grande área superior do Word, que contém todos os seus comandos, organizados na forma de ferramentas de fácil acesso.



*Figura 6.2 – Faixa de Opções do Word 2010.*

Só um detalhe acerca da Faixa de Opções: ela “matou” a interface anterior, que apresentava “Barras de Menus” e “Barras de Ferramentas”. Ou seja, nos Word 2007 e 2010, não há mais menus, nem barras de ferramentas... Só há a **Faixa de Opções**!

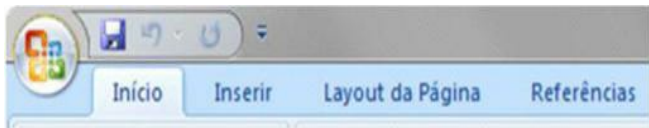
## Guias

A Faixa de Opções, por sua vez, está dividida em guias (ou abas): Arquivo, Página Inicial, Inserir, Layout de Página, Referências, Correspondências, Revisão e Exibição são as guias do Word 2010.



*Figura 6.3 – Algumas guias da Faixa de Opções do Word 2010.*

Só uma pequena diferença entre o Word 2010 e o Word 2007: A guia “Arquivo” não existe no Word 2007... Lá, o botão Office (o botão redondo com o símbolo do Office) é que faz o papel da guia Arquivo. Além disso, no Office 2007, a guia “**Página Inicial**” tem outro nome: “**Início**”, apenas!



*Figura 6.4 – Algumas guias da Faixa de Opções do Word 2007.*

## Grupos

As guias da Faixa de Opções são divididas em Grupos de Ferramentas, ou, simplesmente, grupos. Cada grupo é uma reunião de ferramentas específicas para um determinado fim.

Na figura a seguir, é possível ver os grupos *Área de Transferência*, *Fonte* e *Parágrafo*.



Figura 6.5 – Alguns grupos da guia Página Inicial.

Note que na parte inferior de cada grupo está seu nome. Repare, também, que, à direita do nome, há um pequeno ícone. Este ícone (ou botão) dá acesso a uma janela com mais opções acerca daquele grupo.

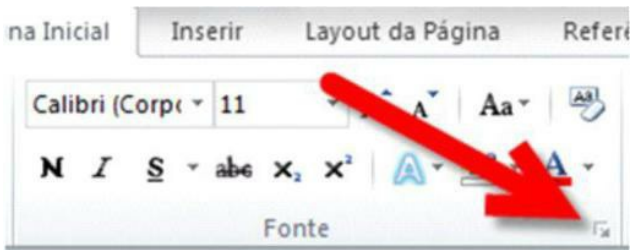


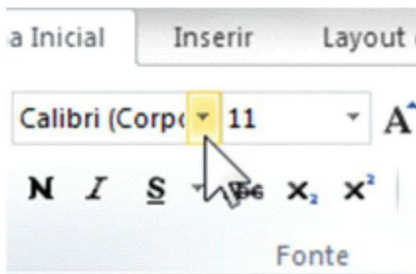
Figura 6.6 – Botão que dá acesso à janela Fonte.

Vamos analisar as janelas abertas por cada grupo mais adiante, quando analisarmos cada um dos grupos da Faixa de Opções do Word 2010.

## Ferramentas

Finalmente, a cada botão, ou comando, disponível nos diversos grupos nas guias da Faixa de Opções, damos o nome de Ferramentas (ou botões, ou comandos, sei lá... Você escolhe!).

As ferramentas não precisam ser, necessariamente, *botões* (do tipo que se clica), como os comandos *Negrito* e *Itálico*... Podem ser de outros formatos, também, como *Drop Down* (caixas de listagem ou caixas de combinação), como os comandos de *Tipo da Fonte* e *Tamanho da Fonte*, onde há vários itens, apresentados em listas, para escolher.



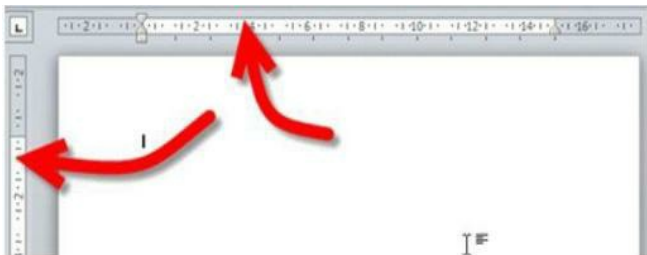
*Figura 6.7 – Ferramentas (Detalhe no Tipo da Fonte).*

Vamos dar mais detalhes das Ferramentas, analisando inclusive suas teclas de atalho, mais adiante. Não se preocupe! Tem muita coisa vindo aí!

#### 6.1.1.2. Régua

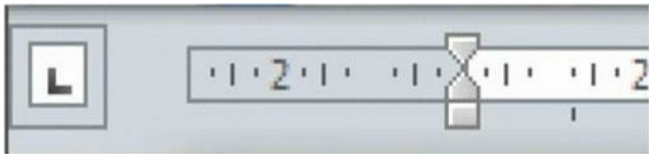
As régua apresentam uma forma simples e rápida de medir a página e as informações do documento a ser digitado.

A régua horizontal mostra (em cinza) as margens esquerda e direita da página (partes da página que não poderão ser usadas pelo texto), e a régua vertical mostra (também em cinza) as margens superior e inferior da página.



*Figura 6.8 – Régua (horizontal e vertical) do Word.*

A régua horizontal também mostra alguns componentes bem interessantes, mostrados em detalhes na figura seguinte:



*Figura 6.9 – Detalhe da régua horizontal (extremidade esquerda).*

### **Controle do Recuo Especial – Primeira Linha**

A setinha localizada na parte superior da régua (ou seja, o “triângulo de cabeça para baixo”) controla o recuo da primeira linha do parágrafo. Ou seja, ao arrastar tal setinha, apenas a primeira linha do parágrafo selecionado apresentará o recuo (afastamento em relação à margem da página).

### **Controle do Recuo Especial – Deslocamento**

A setinha que fica logo abaixo (“triângulo de ponta para cima”) serve para determinar o recuo deslocado (ou deslocamento). Consiste, tão somente, no afastamento das outras linhas do parágrafo (com exceção da primeira linha). Ou seja, se o usuário arrastar essa setinha (em vez do quadradinho), ele irá causar afastamento em todas as linhas do parágrafo, exceto na primeira.

Quando esse triângulo é arrastado, ele leva consigo o quadradinho abaixo dele, mas não carrega a setinha de cima (controle de recuo da primeira linha).

### **Controle do Recuo do Parágrafo – à Esquerda**

Por fim, abaixo de todos há um quadradinho que, quando arrastado, vai ajustar o recuo do parágrafo (o recuo de todas as linhas do parágrafo, sem exceção). Quando esse quadrado é arrastado, ele move as duas setinhas!

### **Selecionador de Marca de Tabulação**

Na lateral esquerda da régua, há um campo quadrado que está mostrando um “L” (uma letra “L” aparentemente). Esse campo controla a escolha de marcas de tabulação (calma, veremos mais adiante o que é isso!).

### **Controle do Recuo do Parágrafo – à Direita**

Na extremidade direita da régua horizontal, que não pode ser vista na figura anterior, existe um solitário triângulo que aponta para cima. Ao arrastá-lo, você estará deslocando o recuo à direita.

Vamos falar mais de recuo em breve... Não se preocupe!

#### **6.1.1.3. Barra de status**

É a barra horizontal, localizada na base da tela do Word que apresenta várias informações a respeito do estado da janela do programa. Consultar a barra de status, especialmente em concursos públicos que usam fotografias, sempre foi de grande ajuda para os candidatos!



*Figura 6.10 – Parte da barra de status do Word.*

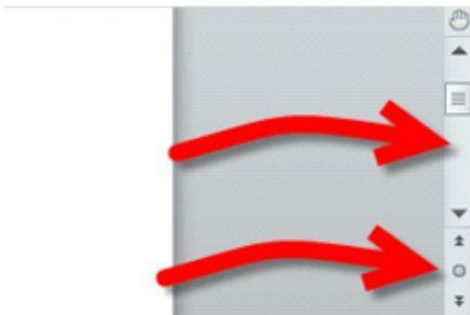
#### **6.1.1.4. Barra de rolagem**

Oferece recursos para rolar o conteúdo da tela de modo que se possa visualizá-lo completamente. Há barras de rolagem sempre que o conteúdo total do documento (páginas, textos, figuras) não puder ser apresentado em uma única tela.

“Mas só existe a barra de rolagem vertical no Word, não é, João?”

Não, caro leitor! Também há a barra de rolagem horizontal, quando o documento estiver sendo visualizado de tal forma que não dê para mostrar todo o conteúdo lateralmente.

Além de permitir realizar a rolagem simples da tela, há comandos abaixo da rolagem vertical que permitem a localização mais rápida de certos objetos no texto (veremos esses recursos mais adiante).



*Figura 6.11 – Barra de rolagem vertical e os botões de navegação.*

#### **6.1.2. Digitando no Microsoft Word**

Para digitar no Word é muito simples: basta posicionar o *ponto de inserção* (aquela barrinha

que fica piscando e que muitos chamam de cursor) no local desejado e começar a digitar. As palavras que não couberem em uma determinada linha de texto serão imediatamente jogadas para a próxima linha.

Para posicionar o ponto de inserção em qualquer local do texto usando o mouse, basta clicar no local desejado.

Também podemos fazer uso de algumas teclas para posicionar o ponto de inserção no local correto. As teclas mais usadas no Word e suas funções são:

- a. As teclas de **seta para esquerda** e **seta para a direita** servem para mover o ponto de inserção um caractere para a direção a que apontam;
- b. As teclas de **seta para cima** e **seta para baixo** permitem que o ponto de inserção suba uma linha ou desça uma linha respectivamente;
- c. A tecla **HOME** faz o ponto de inserção se posicionar no início da linha atual;
- d. A tecla **END** faz o ponto de inserção se posicionar no final da linha atual;
- e. A tecla **ENTER** faz o Word realizar uma quebra de parágrafo (ou seja, o parágrafo atual será encerrado e o ponto de inserção será posicionado no início de um novo parágrafo, numa linha abaixo da atual);
- f. A tecla **BACKSPACE** apaga o caractere que estiver imediatamente à esquerda (ou seja, antes) do ponto de inserção;
- g. A tecla **DELETE** (ou DEL) apaga o caractere que estiver posicionado à direita (ou seja, depois) do ponto de inserção;
- h. As teclas **PAGE UP** e **PAGE DOWN** não fazem exatamente o que dizem. Essas teclas não servem para passar de uma página para outra, mas para fazer a tela (parte visível da página) rolar para cima (Page Up) ou para baixo (Page Down);
- i. A tecla **TAB** serve para inserir um caractere de tabulação (um “símbolo” que diz ao texto para saltar para a próxima marca de tabulação – ou seja, simplesmente, o ponto de inserção dá um “salto” à frente...).
- j. A tecla **SHIFT** deve ser mantida pressionada para acessar as funções secundárias das teclas e as letras maiúsculas. Essa tecla também é utilizada em processos de seleção de trechos do texto (veremos adiante);
- k. A tecla **CAPS LOCK** serve para travar a caixa alta das letras, ou seja, basta pressioná-la uma vez e todas as letras do teclado serão escritas em maiúsculas; pressionando-a outra vez, as letras voltam a ser escritas em minúsculas.

Note que isso só serve para textos que ainda serão digitados. Se você quiser transformar um trecho já digitado em maiúsculas, de volta para minúsculas, não é o CAPS LOCK que faz isso!

Há ainda algumas funções de movimento quando combinamos a tecla **CTRL** com algumas das que vimos anteriormente, veja a seguir:

- a. Se o usuário mantiver a tecla **CTRL** pressionada e acionar as setas **esquerda** ou **direita**, o ponto de inserção só saltará entre os inícios das palavras (anterior ou posterior, respectivamente);
- b. Se o usuário mantiver a tecla **CTRL** pressionada e acionar as setas **acima** ou **abaixo**, o ponto de inserção só saltará entre os inícios dos parágrafos (anterior ou posterior,



respectivamente);

c. **CTRL + HOME** faz o ponto de inserção se posicionar no início do texto (início da primeira página do arquivo);

d. **CTRL + END** faz o ponto de inserção se posicionar no final do texto (fim do arquivo);

e. **CTRL + DELETE** apaga uma palavra inteira à direita do cursor (ponto de inserção);

f. **CTRL + BACKSPACE** apaga uma palavra inteira à esquerda do cursor (ponto de inserção);

“É só isso, João? Word é só isso?”

Não, amigo leitor! Também temos que saber selecionar trechos de texto, como, por exemplo, nos momentos em que queremos aplicar efeitos aos trechos em questão.

#### 6.1.2.1. Conhecendo o texto

Eis um conjunto de conceitos que pode até parecer “besta” ou “desnecessário”, mas não é! É importante que se saiba o que é cada um deles, ok?

- **Caractere:** cada letra, número, símbolo (até mesmo o espaço) que digitamos em um texto. Por exemplo, a expressão “*Casa Amarela*” tem 12 caracteres (o espaço também conta).

- **Palavra:** conjunto de caracteres (letras e números) que termina com um espaço ou uma pontuação (vírgula, ponto, ponto e vírgula, exclamação etc.). As pontuações e os espaços nunca são considerados parte de palavra nenhuma!

- **Linha:** uma única linha horizontal de texto (como esta).

- **Frase (ou Sentença):** no Word, uma frase é um bloco de texto (vários caracteres) que termina, necessariamente, com ponto, exclamação ou interrogação (reticências, oficialmente, não é fim de frase... Mas da forma como a digitamos – três pontos seguidos – o Word acaba por entender que é fim de frase também).

- **Parágrafo:** é um bloco de texto que termina numa **Marca de Parágrafo** (também chamado **de Caractere de Parágrafo**), que é o símbolo inserido cada vez que digitamos ENTER.

Dá uma olhada na figura a seguir... Vamos entender algumas coisinhas...

Este é o primeiro parágrafo do texto. Aqui começa a  
segunda frase deste primeiro parágrafo. Este  
parágrafo contém três frases.¶

Já comecei a digitar o segundo parágrafo de texto  
de exemplo. Finalmente: chegamos à segunda e  
última frase — é a última mesmo!¶

*Figura 6.12 – Analise o texto aí contido...*

Ele possui 253 caracteres (contando os espaços), 42 palavras, em 2 parágrafos, cada um com 3 linhas de texto. O primeiro parágrafo possui 3 frases, fáceis de localizar pelos “pontos” que as finalizam.

O segundo parágrafo, por sua vez, tem apenas 2 frases. Os sinais de “dois pontos” e “travessão” fazem parte da segunda frase, mas não a finalizam nem delimitam – apenas fazem parte de seu conteúdo.

Você pode ter notado a diferença “visual” no texto acima, caro leitor: apareceram uns “símbolos” estranhos, como uns quadradinhos entre as palavras e um troço esquisito (¶).

Esse é o modo “Mostrar Tudo”, no qual podemos ver os caracteres “não imprimíveis”, como espaços (os quadradinhos entre as palavras), marcas de parágrafo (ou seja, ENTER), marcas de tabulação, entre outros...

Esse **troço estranho** é justamente o ENTER (indica o fim de um parágrafo).

Vamos falar o modo Mostrar Tudo mais adiante. Com mais detalhes!

### 6.1.3. Selecionando trechos do texto

Selecionar é escolher o trecho do texto com o qual se vai trabalhar, para, por exemplo, aplicar certos efeitos a ele. Quando selecionamos um texto, ele fica envolvido por uma “tarja” azul...

## Trecho de **texto selecionado** no Word 2010.

Figura 6.13 – A expressão “texto selecionado” está selecionada nesse texto.

Existem várias formas de selecionar uma série de trechos diferentes, a seguir estão listadas algumas maneiras:

- a. Para selecionar qualquer trecho, com qualquer quantidade de letras ou palavras, deve-se apenas ***clique no início do trecho desejado e arrastar o mouse até o final do mesmo***, que será indicado durante a execução do movimento;
- b. Para selecionar apenas uma palavra, o usuário pode aplicar ***um clique duplo na palavra*** desejada;
- c. Para selecionar apenas um parágrafo, o usuário pode aplicar ***um triplo clique em qualquer palavra*** inserida no parágrafo desejado;
- d. Para selecionar uma frase, mantenha pressionada ***a tecla CTRL*** e clique em qualquer palavra da frase desejada.

Podemos, também, posicionar o mouse na ***margem esquerda da página*** (Isso fará o mouse se transformar em uma setinha branca apontando para a direita). Uma vez posicionado nesse local, as funções dos cliques do mouse serão diferentes:

- e. Se acionarmos ***um único clique***, o Word selecionará apenas a linha do texto para onde nosso ponteiro estiver apontando;
- f. Se acionarmos ***duplo clique***, o Word selecionará o parágrafo;
- g. Se acionarmos ***triplo clique***, o Word selecionará todo o texto do arquivo. Esse comando é equivalente ao comando Selecionar Tudo.
- h. Se ***clícarmos*** enquanto seguramos ***a tecla CTRL***, também selecionamos o texto todo.

Outra maneira de selecionar um trecho específico no Word é usando a tecla SHIFT. Basta manter a tecla SHIFT pressionada e movimentar o ponto de inserção.

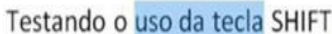
Toda tentativa de movimentar o ponto de inserção (seja com o teclado, seja com o mouse), em conjunto com a tecla SHIFT, fará o Word selecionar o trecho envolvido pelo movimento.

Um exemplo simples: veja a expressão “Testando o uso da tecla SHIFT”. Se o usuário realizar

a seguinte sequência de ações:

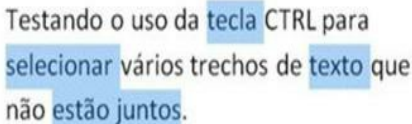
1. Clicar antes a letra “u” de “uso”;
2. Segurar a tecla SHIFT;
3. Finalmente, clicar depois da letra “a” de “tecla”.

O resultado é o seguinte:



*Figura 6.14 – Exemplo de seleção com SHIFT.*

Também é possível selecionar diversos trechos diferentes simultaneamente no Word. Para selecionar várias partes do texto, apenas selecione o primeiro trecho e, segurando a tecla CTRL, selecione os demais trechos (não importa como você as seleciona, se com arrasto, duplo clique, triplo clique etc.).



*Figura 6.15 – Selecionando múltiplos trechos de texto.*

“Mas, João, o CTRL serve para três coisas? Selecionar frases, selecionar o texto todo, selecionar trechos múltiplos. Como diferenciar as três?”

Fácil, nobre leitor. É só se perguntar isto: “Ao segurar a tecla CTRL, havia algo selecionado antes?” Se a resposta for afirmativa, então o CTRL servirá para selecionar mais de um trecho. (É só ter algo selecionado antes para que o Word identifique que o CTRL assumirá a posição de “selecionador múltiplo”).

Caso, na hora de pressionar o CTRL, não haja nada previamente selecionado, o CTRL vai assumir sua posição de “selecionador de frase” (se o clique for dado em qualquer parte do texto) ou de “selecionador do texto todo” (se o clique for dado em qualquer parte da margem esquerda da página).

#### 6.1.3.1. Alguns detalhes sobre trechos selecionados

- **Lembrete 1:** Enquanto um trecho está selecionado, quando pressionamos as teclas Delete ou Backspace, todo o trecho selecionado é imediatamente apagado. (Isso acontece, também, quando há mais de um trecho selecionado.)
- **Lembrete 2:** Se um trecho está selecionado e pressionamos uma tecla para inserir um caractere qualquer (uma letra, por exemplo), o trecho selecionado é imediatamente substituído pelo caractere digitado. (Cabe aqui lembrar que espaço, ENTER e TAB também são caracteres.)
- **Lembrete 3:** Se o ponto de inserção estiver entre duas letras de uma palavra, qualquer efeito de formatação de caracteres será aplicado a toda a palavra. (Portanto, não é necessário selecionar a palavra toda quando se quer aplicar nela efeitos de letra.)

A mesma ideia serve para os efeitos de parágrafo. Ou seja, se você quer, por exemplo, aplicar o alinhamento justificado a um parágrafo inteiro, não é necessário selecioná-lo por inteiro (sei, sei... você fazia isso o tempo todo, não é?). Basta colocar o ponto de inserção em qualquer local daquele parágrafo.

“João, você citou ‘efeitos de letra’ e ‘efeitos de parágrafo’. O que são eles?”

#### 6.1.3.2. Efeitos de caractere (fonte) *versus* efeitos de parágrafo

Aproveitando sua pergunta, leitor, vamos classificar os efeitos que o Word pode aplicar no texto:

- **Efeitos de caractere (ou efeitos de letra, ou efeitos de fonte):** são alguns efeitos que o Word pode aplicar diretamente sobre os caracteres (letras, números e símbolos que digitamos no texto). Negrito, itálico e sublinhado são exemplos de efeitos de fonte, pois pode-se aplicá-los, se desejado, a cada letra separadamente.

Também são exemplos de efeitos aplicáveis às fontes: tipo da fonte, tamanho da fonte, cor da fonte, subscrito, sobrescrito, tachado e tachado duplo, baixo e alto relevo, versalete e sombra etc. Não por acaso, esses efeitos se encontram no grupo Fonte.

- **Efeitos de parágrafo:** são recursos que se aplicam aos parágrafos (ou seja, são efeitos inerentes à entidade parágrafo – bloco de texto). Ou seja, tais recursos não apresentam, num mesmo parágrafo, duas formatações diferentes.

Entre os efeitos classificados como efeitos aplicáveis a parágrafos estão o alinhamento do parágrafo (esquerdo, direito, justificado e centralizado), marcadores e numeração, recuos,

afastamento de linha, afastamento antes e depois do parágrafo.

## 6.2. Principais comandos e recursos do Word

O Word tem uma variada coleção de comandos que ajudam o usuário em sua tarefa de criar e editar documentos de texto profissionais. A grande parte desses comandos não é explorada em concursos públicos, mas os mais comuns sempre estão presentes.

Segue uma listagem, dividida pela posição nas guias da Faixa de Opções do programa. Vamos começar pela guia Página Inicial, deixando a guia Arquivo (ou botão Office, no Word 2007) para depois, ok?

### 6.2.1. Guia Página Inicial

A guia Página Inicial (ou *Início*, no Word 2007) contém os mais comuns comandos do programa. Sem dúvida alguma, a maior probabilidade de alguma ferramenta ser exigida em prova é advinda desta guia (claro que podem cair comandos de qualquer uma, é verdade...):

#### 6.2.1.1. Grupo Área de Transferência

Este grupo de comandos contém os principais recursos para recortar/copiar e colar objetos.

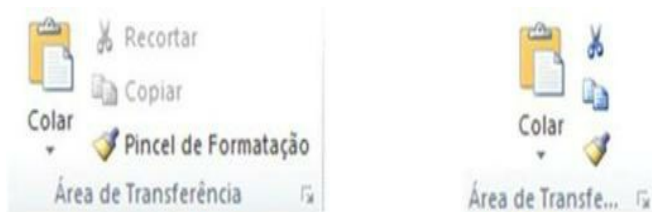


Figura 6.16 – Grupo Área de Transferência (duas versões).

Na figura acima, podemos ver duas “formas” de exibir o grupo Área de Transferência: isso depende somente da largura da janela do Word utilizada. Quanto mais “estreita”, menos informações serão mostradas (como na imagem acima à direita).

Segue a lista dos comandos existentes neste grupo:

#### Colar

Este comando insere, no texto, o último objeto (ou trecho de texto) que havia sido copiado ou recortado recentemente. Para acioná-lo via teclado, basta usar a tecla de atalho **CTRL + V**.

Ao clicar diretamente no botão Colar, será realizada a “colagem” do objeto que estiver na área de transferência (ou seja, o último objeto recentemente copiado ou recortado), mas ao

clicar na setinha abaixo deste botão, outras opções para colagem são apresentadas:



Figura 6.17 – Opções de Colagem.

É possível colar somente texto (caso se tenha copiado texto + figuras), é possível colar texto já se adaptando à formatação do destino (desconsiderando como o texto estava formatado anteriormente), entre outras coisas...

Esse menu também dá acesso à opção **Colar Especial**, que abre uma janela com muitas opções de colagem. A tecla de atalho para o Colar Especial é **CTRL + ALT + V**.

A opção Definir Colagem Padrão permite que você indique qual o tipo normal de colagem que você deseja (ou seja, tornará padrão, a ser obedecido imediatamente nas próximas vezes que você colar o tipo de colagem que quiser!).

## Recortar

Este comando **envia o objeto** (ou trecho de texto) selecionado para a área de transferência. De lá, ele pode ser colado quantas vezes quiser! Sua tecla de atalho é **CTRL + X**.

Só note, por favor, que o objeto é retirado do local original quando é recortado! Ou seja, ele é **MOVIDO**.

Mas, uma coisa interessante: se você recortar um objeto (ou trecho) qualquer e não o cola, ele simplesmente “some” – ou seja, ele é simplesmente apagado.

## Copiar

O comando copiar envia uma cópia do objeto (ou trecho de texto) selecionado para a área de transferência. Note: o objeto (ou trecho) não sai de onde está, pois apenas uma cópia dele é enviada! A tecla de atalho é **CTRL + C**.

## Pincel de Formatação

Este comando copia o formato de fonte (cores, tipo, tamanho) e parágrafo (estilos, alinhamento, afastamento de linha, recuo) de um trecho de texto para outro trecho. Sua tecla de

atalho (para copiar o efeito de formatação do trecho original) é **CTRL + SHIFT + C**, e depois, para aplicar o efeito no trecho de destino, usa-se **CTRL + SHIFT + V**.

### Botão de Controle do Grupo Área de Transferência

O botão que acompanha este grupo (que fica à direita do nome do grupo) dá acesso ao recurso Área de Transferência do Office 2010.

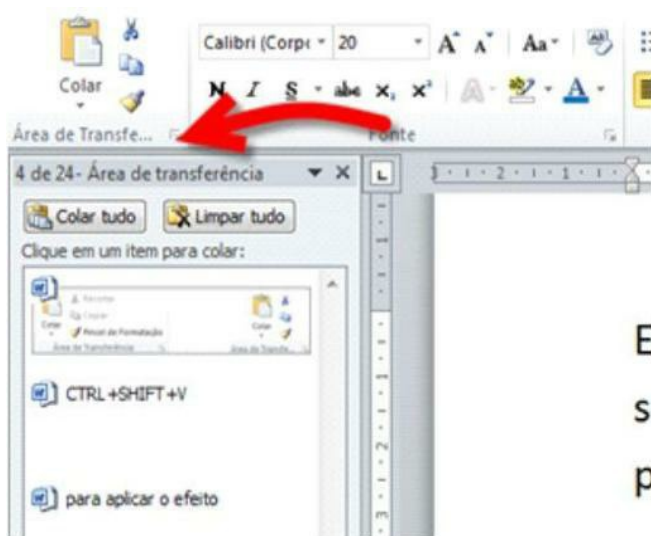


Figura 6.18 – Área de Transferência Aberta.

A Área de Transferência do Office 2010 consegue armazenar até 24 objetos simultaneamente. Você pode copiar ou recortar diversos trechos diferentes (a cada comando de cópia ou recorte, o item é “coletado” para dentro da Área de Transferência) e, então, colar qualquer um deles, onde o ponto de inserção estiver!

#### 6.2.1.2. Grupo Fonte



Os comandos contidos neste grupo dizem respeito às operações realizáveis com as letras (fontes), como efeitos de negrito, itálico etc.



Figura 6.19 – Grupo Fonte.

### Tipo da Fonte

Este controle permite escolher o tipo da fonte (letra) que será aplicado ao texto selecionado. Na figura acima, é a caixa de listagem que indica “Calibri (Corpo)”.

### Tamanho da Fonte

Este comando determina (selecionando em uma lista) o tamanho (em pontos) que a fonte (letra) vai ter. Na figura anterior, é a caixa de listagem que apresenta o número 11.

### Aumentar Fonte e Diminuir Fonte

Esses dois botões (“A” grande com a setinha pra cima, e “A” pequeno com a setinha para baixo) fazem o mesmo que o controle anterior, só que apenas usando cliques (cada clique faz o aumento – ou diminuição – da fonte em alguns pontos). É possível fazer o mesmo via tecla de atalho: **CTRL** + > (para aumentar) e **CTRL** + < (para diminuir).

Um detalhe, apenas: como os símbolos “>” (maior que) e “<” (menor que) são conseguidos na segunda função das teclas “.” (ponto) e “,” (vírgula), respectivamente, podemos ver alguma banca escrevendo tais teclas de atalho da seguinte forma:

**CTRL** + **SHIFT** + , (que é o mesmo que CTRL + >): aumenta o tamanho da fonte.

**CTRL** + **SHIFT** + . (igual a CTRL + <): para diminuir o tamanho da fonte.

### Maiúsculas e Minúsculas

Este simples botão (com ícone “Aa”) abre uma listagem que apresenta opções para alterar a caixa do texto selecionado, oferecendo as seguintes alternativas: MAIÚSCULAS; minúsculas; Colocar Cada Palavra Em Maiúsculas; aLTERNAR mAIÚSC./mINÚSC.; e Primeira letra da sentença em maiúscula. A combinação de teclas que aciona este comando é **SHIFT** + **F3**.

### Limpar Formatação

Este botão (“Aa” com uma borrachinha) retira todos os efeitos de formatação de um trecho selecionado (negrito, tamanho, efeitos de parágrafo, cores etc.), deixando-o no estilo Normal. Não se preocupe que falaremos sobre Estilos logo a seguir!

### Negrito, Itálico e Sublinhado

Estes três comandos (muito conhecidos, por sinal) aplicam efeitos distintos no trecho selecionado. Negrito (botão com o “N”) deixa o **texto mais encorpado (letras mais grossas)**, Itálico (botão com o “I”) deixa o *texto levemente inclinado para a direita* e Sublinhado (botão com o “S”) apresenta uma linha abaixo do trecho onde se aplica o efeito.

Perceba, apenas, que o sublinhado apresenta uma setinha à direita do botão, que permite a escolha de alguns tipos de sublinhados especiais.

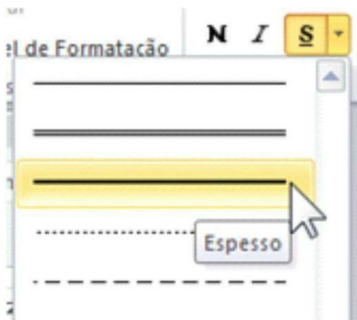


Figura 6.20 – Opções de sublinhado – clique na setinha do botão.

As teclas de atalho para os três comandos são mais do que conhecidas (você tem que sabê-las!) : **CTRL + N** (Negrito); **CTRL + I** (Itálico); **CTRL + S** (Sublinhado).

### Tachado

É o botão do “abc” cortado por uma linha! Este comando desenha uma linha no meio do trecho selecionado onde se aplica o efeito.

### Subscrito e Sobrescrito

Estes dois comandos alteram a linha base dos trechos selecionados, apresentando-os um pouco abaixo ou um pouco acima (respectivamente) da linha base do texto original. O efeito resultante é:

H<sub>2</sub>O (2 está subscrito);

$4^2 = 16$  (agora o 2 está sobrescrito);

As teclas de atalho para esses comandos são: CTRL + = (subscrito); e CTRL + SHIFT + = (sobrescrito).

Pode ser que eles “anunciem” o sobrescrito como tendo a tecla de atalho CTRL + + (ou seja, segurar CTRL e acionar a tecla “+”), mas a tecla “+” é conseguida com o SHIFT na tecla “=”... Por isso é que se pode descrever qualquer um dos dois “jeitos”.

## Efeitos de Texto

O botão que apresenta uma letra “A” em formato “brilhante”, quase “transcendente” (exagero, né?) é chamado Efeitos de Texto. Este comando oferece formas de “enfeitar” o texto que não se viam (nem se podiam) fazer nas versões anteriores do Word.

Na verdade, esses efeitos até que podiam ser feitos, mas num recurso chamado WordArt (texto especial). Hoje, no Word 2010, esses efeitos podem ser aplicados diretamente dentro do texto, do texto comum que digitamos na página!

Olha uma “palhinha” do que se pode aplicar de Efeitos de Texto:



Figura 6.21 – Efeitos de texto possíveis.

## Cor do Realce do Texto

Este comando (cujo botão é o penúltimo do grupo, parecendo um “ab” sendo marcado por uma caneta marca-texto) cria, justamente, um efeito semelhante ao de um marca-texto (aquelas

canetas hidrográficas com cores “discretas”).

Ao clicar na setinha que acompanha este botão, é possível ter acesso à listagem de cores possíveis de serem aplicadas.

### Cor da Fonte

Este comando, que tem a imagem da letra “A” e abaixo dela um retângulo na cor selecionada, permite alterar a cor das letras do texto.

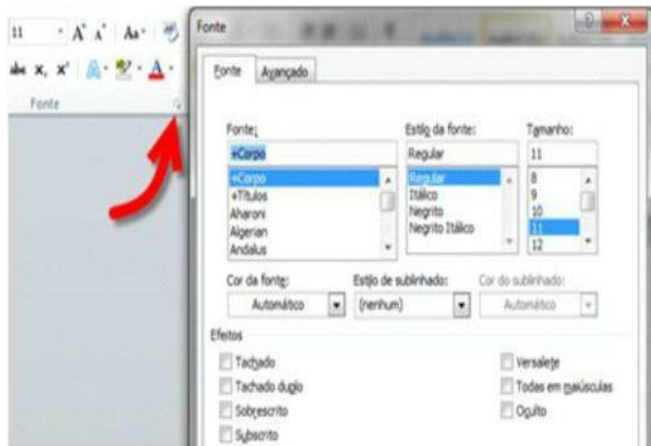
Uma das principais pegadinhas sobre este botão é que ele pode ser confundido com o sublinhado (porque apresenta uma letra e algo abaixo dela que se parece com uma linha). Cuidado!

Outra coisa: ao clicar diretamente no botão Cor da Fonte, será aplicada, ao texto selecionado, a cor que estiver aparecendo no retângulo que fica abaixo da letra “A”, no botão (esse retângulo sempre mostra a última cor escolhida).

Para escolher outra cor, porém, é necessário clicar na setinha que fica à direita deste botão! Ai você terá acesso a uma palheta de cores à sua disposição!

### Botão de Controle do Grupo Fonte

No final do grupo fonte, há a famosa setinha... Esta, especificamente, permite abrir a janela de Opções da Fonte, que oferece uma gama muito maior de efeitos que aqueles apresentados pelas ferramentas do grupo:



### 6.2.1.3. Grupo Parágrafo

Como o nome já diz, este grupo de comandos diz respeito às operações que podem ser realizadas com parágrafos inteiros, como alinhamento, recuos etc.



Figura 6.23 – Grupo Parágrafo.

Para começo de conversa, a setinha que fica no canto inferior direito (ao lado do nome do grupo) serve para abrir a janela de opções de parágrafo (antiga “Formata Parágrafo”, na versão Word 2003 e nas anteriores) – todos os mesmos comandos de antigamente estão nesta janela, como alinhamentos, recuos, espaçamentos de linhas e de parágrafos, linhas órfãs e viúvas etc.

Mas vamos, claro, aos botões do grupo:

#### Marcadores

Esse botão, que apresenta um ícone com três bolinhas e três linhas horizontais, serve para ligar ou desligar os marcadores (“símbolos” simples que indicam o início de um parágrafo no texto – servem para enfeitar a margem esquerda do parágrafo).

Um clique no botão, em si, aciona, ou desliga, o marcador atual (aquele último formato que o usuário escolheu). Um clique na setinha ao lado deste botão permitirá escolher entre vários tipos de marcadores existentes (bolas, quadrados, setas etc.).

#### Numeração

O comando Numeração, cujo ícone traz a imagem de 1, 2, 3, e três linhas horizontais, serve para ligar ou desligar a numeração (um número que vai se incrementando a cada novo parágrafo); novamente, a setinha à direita permite escolher opções de numeração (tipo, tamanho etc.);

#### Lista de Vários Níveis

Este botão, apresentado na forma de 1, a, i e três linhas que vão diminuindo de tamanho, permite ligar ou desligar a numeração de vários níveis (numeração de tópicos/subtópicos); mais uma vez, a setinha à direita permite que se configurem mais opções deste recurso;

## Diminuir Recuo e Aumentar Recuo

Estas ferramentas alteram o recuo do parágrafo selecionado (recuo é o afastamento do texto em relação à margem da página); No caso, estes botões afetam apenas o recuo à esquerda (ou seja, o afastamento do início do texto em relação à margem esquerda da página).

Podemos utilizar também teclas de atalho para acionar os comandos:

**CTRL + M** para aumentar o recuo; e

**CTRL + SHIFT + M** para diminuir o recuo.

## Classificar

O comando Classificar (um botão contendo “A” e “Z” e uma seta para baixo) abre uma janela que permite ordenar os parágrafos do texto (ou linha de uma tabela) de acordo com os critérios de ordem alfabética (para texto) ou numérica (tanto na ordem crescente como decrescente).

## Mostrar Tudo

Este botão (cuja imagem remete à Marca de Parágrafo – o “troço esquisito” já visto anteriormente) permite que o usuário veja (ou oculte novamente) os caracteres não-imprimíveis (como ENTER, TAB, ESPAÇOS e QUEBRAS DE PÁGINA, entre outros). Este recurso pode ser ligado ou desligado por esse botão.

A tecla de atalho usada para este fim é **CTRL + \*** (o asterisco está normalmente em cima da tecla 8, no teclado alfabético, portanto, pode-se dizer, também, **CTRL + SHIFT + 8**).

## Alinhamento de Parágrafo

São os quatro primeiros botões da segunda linha do grupo. São justamente aqueles que têm várias linhas horizontais arrumadas de formas diferentes. São eles, em ordem:

- **Alinhar Texto à Esquerda (CTRL + Q):** alinha o texto na margem esquerda da página, apenas, sem garantir o alinhamento na margem direita;
- **Centralizar (CTRL + E):** alinha o texto no centro da página. As margens esquerda e direita ficam sem alinhamento;
- **Alinhar Texto à Direita (CTRL + G):** alinha o texto na margem direita da página, não garantindo que ele ficará alinhado à margem esquerda.
- **Justificar (CTRL + J):** alinha o parágrafo tanto na margem esquerda quanto na direita, deixando uma sensação de “retângulo” no texto. Para fazer isso, esse recurso aumenta os espaços entre as palavras, causando, algumas vezes, efeitos bem desagradáveis de “lacunas” no meio do texto (mas isso é uma “frescura” minha, só citei... não quer dizer que vão cobrar isso!).

## Espaçamento de Linha e Parágrafo

Este botão, que tem o formato de duas setinhas (uma para cima e outra para baixo), ao lado de algumas linhas horizontais, abre, necessariamente, um pequeno menu que apresentará opções de espaçamento entre linhas (espaço vertical, medido em pontos, entre uma linha e outra do parágrafo) e o espaçamento antes e depois de cada parágrafo (espaço especial entre um parágrafo e outro).

Note que, apesar de aparentar, este botão não tem “setinha” à direita! Ou seja, a setinha, em si, faz parte do próprio botão, já que ele sempre abrirá o menu de opções... (diferentemente de outros casos, como “Marcadores”, “Numeração” e “Cor da fonte”, por exemplo, em que clicar no botão é diferente de clicar na setinha).

Podemos usar algumas teclas de atalho para o recurso de Espaçamento de Linhas: **CTRL + 1** (ajustar para espaçamento simples); **CTRL + 2** (ajustar para espaçamento duplo); **CTRL + 5** (ajusta para espaçamento 1,5).

### Sombreamento

O botão do “balde de tinta” serve para colorir o plano de fundo de um trecho de texto (dar cor à área que fica atrás de um trecho de texto).

Este botão tem duas partes e você já deve ter deduzido: se clicarmos diretamente no botão (balde de tinta), será aplicada a cor de fundo já selecionada ao trecho que receber o efeito. Se clicarmos na setinha ao lado dele, poderemos escolher entre várias opções de cores.

### Bordas

Este botão (que parece uma “janela” com linhas tracejadas) serve para inserir (ou retirar) bordas (linhas) ao redor do trecho selecionado. A setinha à direita permite a escolha do tipo e espessura, bem como a localização das bordas a serem inseridas/retiradas.

Só um aviso: quer ver essas ferramentas em ação? Mexa nelas! Abra o Word 2010 e vá em frente! Usar cada uma delas vai fazer com que você tenha muito mais familiaridade com essas ferramentas, permitindo que você identifique mais facilmente o que se pede delas nas provas!

#### 6.2.1.4. Grupo Estilo

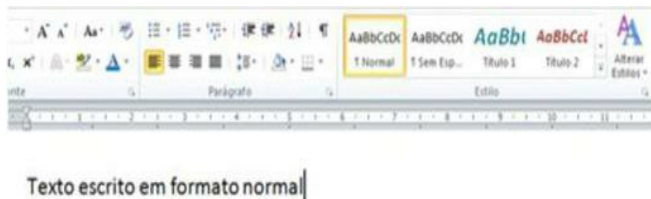
O Grupo Estilo permite que se apliquem, rapidamente, estilos de texto e de parágrafo ao documento.

Estilos são conjuntos de definições de formato que se podem aplicar várias vezes num documento, como por exemplo, o estilo “Título” que pode ser configurado para ser: fonte Arial, tamanho 18, negrito, itálico, vermelho. Uma vez indicando esta definição, sempre que se aplicar o estilo “Título” a um trecho, ele ficará exatamente deste jeito.

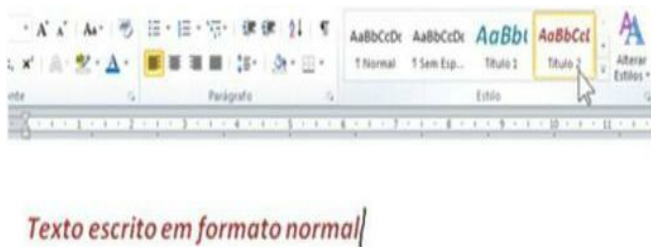


Figura 6.24 – Grupo Estilo.

Um tipo de comando novo é apresentado aqui neste grupo: a “Galeria”. Galeria é um tipo de comando que apresenta várias opções de formatação (“formato”), onde é possível apenas “apontar” o mouse para o efeito desejado e ele é mostrado diretamente no texto... Ao tirar o mouse do efeito na galeria, o trecho de texto volta a mostrar-se como realmente é.



*Figura 6.25 – Texto em estilo Normal.*



*Figura 6.26 – Mouse “pousando” sobre o estilo Título 2, na galeria.*

Ou seja, não é necessário aplicar um efeito para saber como ele ficará no texto, basta colocar o ponteiro do mouse sobre o estilo desejado e o texto selecionado ficará exatamente naquele



formato... Se você retirar o mouse do estilo, o trecho de texto automaticamente voltará à sua normalidade.

Além de aplicar estilos pré-configurados, podem-se criar novos estilos ou alterar os já existentes. Além disso, é possível determinar, através do botão Alterar Estilos, os conjuntos de cores e fontes para um grupo de estilos (assim, pode-se alterar completamente a “cara” do documento com apenas alguns cliques).

Outra coisa para a qual os estilos são muito importantes é a classificação do documento em tópicos e subtópicos, criando uma verdadeira “estrutura” no documento que será de vital importância para a criação de índices automáticos (como sumários).

Para “entender” mais um pouco sobre estilos do Word, assista ao vídeo sobre este assunto no hotsite deste livro na Editora Campus/Elsevier ([www.elsevier.com.br](http://www.elsevier.com.br)) ou assista ao curso de Word 2010 do Eu Vou Passar ([www.euvoupassar.com.br](http://www.euvoupassar.com.br)).

#### 6.2.1.5. Grupo Edição

Este grupinho quase imperceptível à direita da guia Página Inicial contém comandos que antes existiam no menu Editar (nas versões anteriores do Word). Vamos dar uma olhada nos três comandos deste grupo...

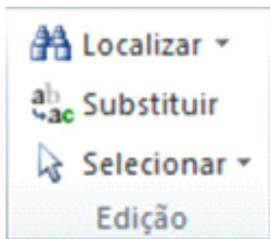


Figura 6.27 – Grupo Edição (bem babaquinha...).

##### Localizar

O botão Localizar (símbolo do binóculo) pode ser clicado diretamente ou na sua setinha à direita (são dois comportamentos diferentes).

Clicando diretamente no botão Localizar, abre-se um campo de pesquisa num painel à esquerda da página do documento. Uma vez digitando um critério de pesquisa, todas as ocorrências deste critério serão marcadas com uma tarja amarela.

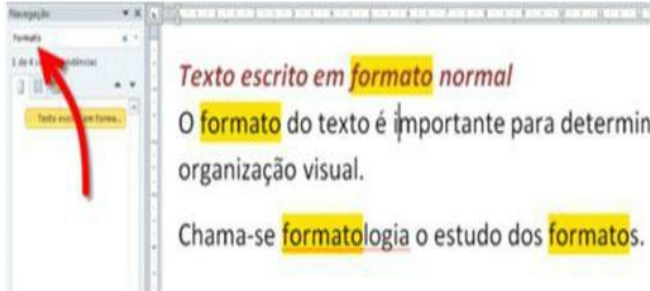


Figura 6.28 – Localizando o trecho “formato”.

Clicando na setinha ao lado, pode-se escolher entre as opções:

- **Localizar (idêntica a clicar no botão):** pode-se acionar, também, pela tecla de atalho **CTRL + L**.
- **Localização Avançada:** abre uma janela para Localização de trechos de texto com muito mais opções (esta era a janela aberta antigamente pelo comando Editar/Localizar, nas versões do Word anteriores – 2003 e mais antigas).

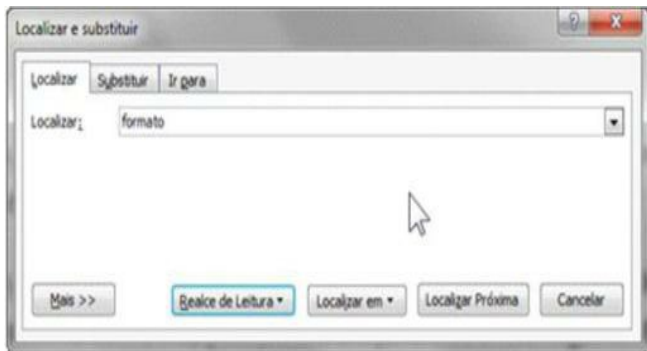


Figura 6.29 – Janela de Localização Avançada.

- **Ir Para:** abre uma janela para permitir que o usuário salte para determinados pontos no documento (“ir para a página 8”, ou “ir para a tabela 15”, ou ainda “ir para a figura 34”). A tecla de atalho para este recurso é **F5**, simplesmente.

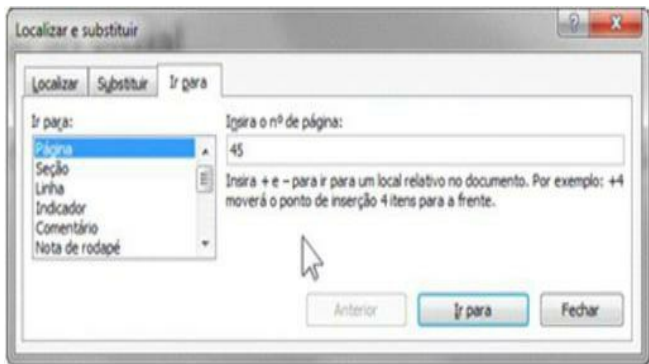


Figura 6.30 – Janela Ir Para.

Não sei se você percebeu, leitor, mas a janela do Ir Para é a mesma do comando Localização Avançada (só muda a guia na parte superior). Portanto você pode acionar um deles e ir para o outro só trocando lá em cima!

### Substituir

Abre a mesma caixa de diálogo dos dois comandos acima, só que na guia substituir, que permite, além de encontrar trechos, substituí-los por outros trechos.

Você pode, por exemplo, pedir que todas as ocorrências da palavra “formato” sejam substituídas, imediatamente, por “formatação”, ou, sei lá... Qualquer coisa que você queira!

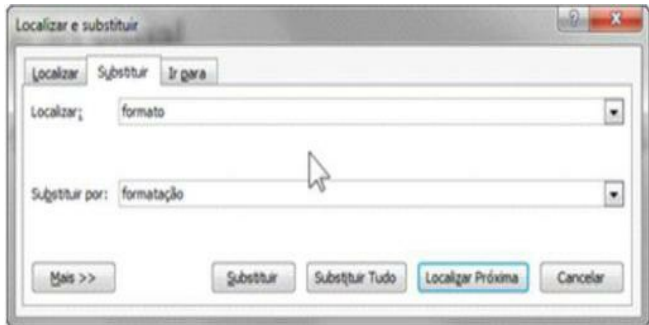


Figura 6.31 – Janela do comando Substituir.

### Selecionar

Este comando, com o símbolo do ponteiro do mouse, necessariamente abre uma listagem com quatro opções:

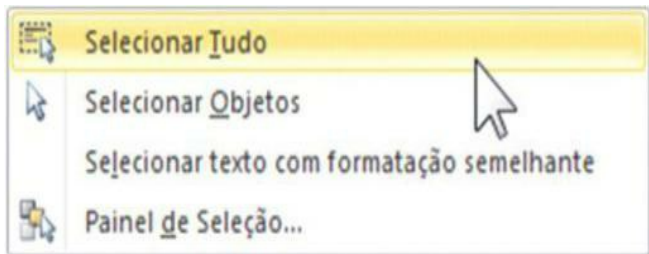


Figura 6.32 – Ao clicar no comando Selecionar, é possível...

- **Selecionar Tudo:** seleciona o texto todo do documento (tecla de atalho: **CTRL + T**);
- **Selecionar Objetos:** transforma o mouse em uma seta branca que somente seleciona objetos (como figuras e autoformas) e não consegue selecionar texto. Um clique novamente nesta ferramenta a desliga, permitindo selecionar texto novamente;

- **Selecionar texto com formatação semelhante:** seleciona todos os trechos de texto do documento que apresentam o mesmo formato (estilo, tipo de fonte, cor da fonte) que o texto que está atualmente selecionado;
- **Painel de Seleção:** abre um painel à direita do documento para permitir a seleção de objetos (figuras, formas) mais facilmente. Este recurso só serve se houver figuras ou formas em seu documento (e se elas estiverem flutuando pelo texto – ou seja, não alinhadas a ele);

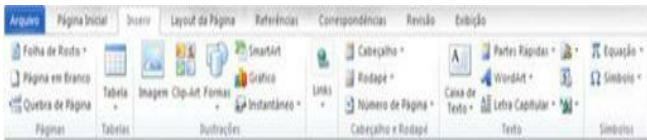
Bem, caro leitor (ou leitora, né?), com isso terminamos a guia Página Inicial... Vamos à guia Inserir agora!

## 6.2.2. Guia Inserir

A guia Inserir traz comandos para colocar objetos e trechos dentro do documento do Word. Ela pode ser vista, a seguir, em dois formatos (um mais “largo”, quando a janela do Word é mais ampla lateralmente, e outro mais “apertado”, quando não se tem tanto espaço na janela):



*Figura 6.33 – Guia Inserir mais “cheinha”.*



*Figura 6.34 – Guia Inserir mais “apertada”.*

Note que os comandos são exatamente os mesmos! Só mudam de “formato” e, mesmo assim, muito pouco (apenas entre botões grandes e pequenos).

### 6.2.2.1. Grupo Páginas

Este grupo possui apenas três comandos bem simples:

#### Folha de Rosto

Este comando, na forma de uma “folha de papel azul”, permite inserir, automaticamente, uma “capa” para o documento (sim, uma “folha de rosto”) que será colocada como página 1,

empurrando as demais páginas para números seguintes. É fantástico como é simples, caro leitor (e meio “desnecessário”).

### Página em Branco

Este comando insere uma página em branco no local onde o ponto de inserção estiver. Caso o ponto de inserção esteja no meio da página 5, por exemplo, será criada uma página 6 (totalmente em branco) e todo o texto que estiver após o ponto de inserção irá para a página 7.

### Quebra de Página

Este botão insere, na posição onde se encontra o ponto de inserção, um caractere especial chamado Quebra de Página, que empurra todo o texto que estiver depois do ponto de inserção para a próxima página. A tecla de atalho é **CTRL + ENTER**.

#### 6.2.2.2. Grupo Tabelas

Esse é muito simples, pois tem apenas uma única ferramenta! Através deste grupo/botão, pode-se inserir uma tabela no documento.

Um único clique neste botão vai abrir uma pequena matriz de exemplo (onde se pode escolher o número de linhas e colunas da tabela a ser criada), além de algumas opções, mostradas e explicadas a seguir:



### Inserir Tabela...

Abre uma janela com algumas opções para a inserção da tabela. Esta janela é mostrada logo a seguir:

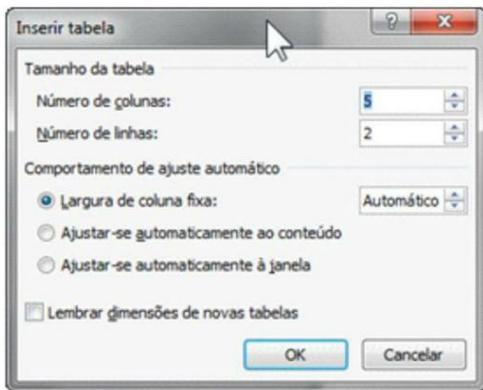


Figura 6.36 – Janela do comando Inserir Tabela.

### Desenhar Tabela

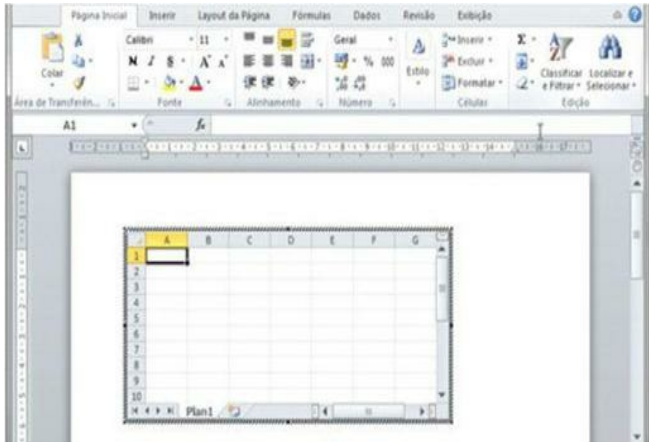
Transforma o ponteiro do mouse em um “lápis” que permite determinar os limites e divisórias da tabela a ser criada por meio de simples “arrastos” do mouse (realmente, como se estivéssemos “desenhando” a tabela em questão).

### Planilha do Excel

Cria, no local onde o ponto de inserção estiver localizado, uma planilha do Excel (que, depois de editada), passará a ser vista, no documento, como um objeto especial (chamado “objeto planilha do Excel”).

Visualmente semelhante a uma tabela, este tipo de objeto não pode ser editado tão livremente quanto uma tabela, mas tem a vantagem de usar todo o poder do Excel.

Para editar um objeto Planilha, deve-se acionar duplo clique nele. Quando isso acontece, o objeto é mostrado como sendo parte do Microsoft Excel, mas dentro do documento do Word.



*Figura 6.37 – Objeto Planilha do Excel dentro do documento do Word.*

### **Tabelas Rápidas**

Apresenta uma lista de modelos predefinidos e já preenchidos de tabelas para serem inseridas. Basta clicar em qualquer um dos modelos apresentados e ele será automaticamente inserido no documento.

Lembre-se, leitor: essas tabelas já vêm preenchidas (com conteúdo) e, claro, podem ser livremente editadas (o conteúdo predeterminado pode ser alterado!).





Figura 6.38 – Menu Tabelas Rápidas.

### 6.2.2.3. Grupo Ilustrações

Traz diversos comandos para a inserção de figuras e outros objetos gráficos no documento do Word.



### **Imagem**

Esse botão permite que se insira no documento uma imagem gravada na forma de um arquivo em seu computador (com as extensões JPG, GIF, WMF, BMP, PNG entre outras).

### **Clip-Art**

Esse botão dá acesso à galeria de clip-arts (figuras que já vêm acompanhando o Office 2010). Se houver acesso à Internet no momento da execução do comando em questão, o Word permitirá acesso às galerias on-line (no site do Office 2010).

O Painele dos Clip-arts não mostra, inicialmente, as figuras disponíveis: ele dá a você um campo para pesquisá-las. Você pode propor, por exemplo, o critério “Eu Vou Passar” na pesquisa e o Word irá trazer todas as imagens de Clip-Art disponíveis para esse “tópico”.

### **Formas**

Esse comando permite a inserção de formas geométricas simples, como círculos, setas, retângulos, balões etc. As formas poderão ser alteradas em uma série de características, como cor, tamanho, efeitos 3D, entre outros.

### **SmartArt**

Essa foi uma novidade do Office 2007 (a versão anterior). Esse recurso permite que se desenhem, no documento, gráficos, fluxogramas, esquemas e diagramas visuais onde se podem inserir textos. Com poucos cliques, pude construir, por exemplo, caro leitor, o esquema abaixo:



*Figura 6.40 – Exemplo de SmartArt (um fluxograma).*

## Gráfico

Essa ferramenta abre a janela que permite a inserção de gráficos com o auxílio do Excel. Sim! Atenção! Apesar de o gráfico ser inserido dentro do Word, as séries de dados (textos e números) a partir dos quais o gráfico foi construído são manipuladas no Excel diretamente.

## Instantâneo

Esta ferramenta permite que se insira uma foto da tela (mais precisamente de qualquer janela aberta e *não minimizada* no computador naquele momento) ou uma foto de um trecho específico da tela (um “recorte da tela”), a ser selecionado com o mouse após a execução do comando.

### 6.2.2.4. Grupo Links

Este pequeno grupo traz recursos de interatividade Web. Com as ferramentas deste grupo, é possível inserir no documento:

#### Hiperlink

Cria um vínculo (apontador) para qualquer recurso acessível (páginas da Internet, outros documentos do Word, outros arquivos de diversos tipos, pontos “indicadores” específicos dentro do documento, endereços e e-mail etc.).

Na verdade, o comando Hiperlink abre uma janela onde se pode descrever o endereço para o

qual o hiperlink apontará.

A tecla de atalho para este comando (ou seja, para abrir a janela) é **CTRL + K**.

### **Indicador**

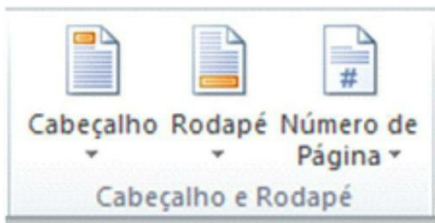
Cria um “nome” para um ponto específico num documento. Quando o usuário cria um indicador, aquele ponto onde o indicador foi criado recebe um nome, que servirá para identificar aquele local específico num hiperlink. Ou seja, quando se cria um indicador, pode-se apontar para aquele indicador usando um hiperlink.

### **Referência Cruzada**

Cria um vínculo (um hiperlink) entre partes de um documento, como por exemplo: num texto que se refira a uma figura – “veja mais como mostrado no gráfico 3.2”. Um clique na referência cruzada levará ao gráfico em questão.

#### **6.2.2.5. Grupo Cabeçalho e Rodapé**

Como o nome já diz, esse grupo apresenta ferramentas que manipulam cabeçalhos e rodapés nos documentos.



*Figura 6.41 – Grupo Cabeçalho e Rodapé.*

### **Cabeçalho**

Essa ferramenta insere conteúdo no cabeçalho do documento (a área de cima da página). Tudo o que você inserir nessa área especial será apresentado em todas as páginas de uma mesma seção (e, se você quiser, em todas as páginas do documento inteiro).

Ao clicar no botão Cabeçalho, será apresentada uma listagem de opções de modelos de cabeçalhos e alguns outros comandos como: Editar Cabeçalho e Remover Cabeçalho.

### **Rodapé**

Idêntico ao cabeçalho (com a diferença de referir-se à parte inferior da página), o rodapé repete-se em todas as páginas de uma mesma seção (ou do documento inteiro, se o usuário assim quiser).

Novamente, um clique nesta ferramenta dará acesso a uma lista de opções de modelos predefinidos de rodapés. Editar e Remover rodapés também fazem parte das opções possíveis.

### Número de Página

Insere um número automático para as páginas do documento, colocando tal informação na posição que o usuário determinar. Ao clicar no botão, surge um menu contendo as opções:

- **Início da Página:** insere números de páginas no cabeçalho do documento;
- **Fim da Página:** insere números de páginas no rodapé do documento;
- **Margens da Página:** insere números de páginas nas laterais do documento (margens esquerda ou direita);
- **Posição Atual:** insere o número da página atual na posição em que o ponto de inserção estiver (não se repete em todas as páginas, apenas se for colocado no cabeçalho ou no rodapé da página).

Ainda se encontram as opções **Formatar Números de Página** e **Remover Números de Página**.

#### 6.2.2.6. Grupo Texto

Esse grupo trabalha com uma série de ferramentas que vão permitir a inserção de recursos de texto no documento em questão.



Figura 6.42 – Grupo Texto.

### Caixa de Texto

Insere uma caixa de texto, que é uma “moldura” dentro da qual se pode escrever texto. Normalmente criamos caixas de texto para colocar trechos de texto em posições independentes do fluxo natural do texto no documento.

### Partes Rápidas

Insere conteúdo de texto reutilizável (será possível inserir várias vezes os mesmos conteúdos), como as propriedades do documento (assunto, autor, empresa, telefone etc.) e campos automáticos (são áreas de conteúdo mutável, alterável pelo próprio Word, como número da página, data atual, hora atual, índices...).

Não esqueça, ok? Campos (variáveis de texto) são inseridos por meio da opção Partes Rápidas,

do grupo Texto!

## **WordArt**

Insere um texto decorativo (cheio de frescuras)... São pequenos trechos (algumas poucas palavras, normalmente) que podem receber inúmeros efeitos interessantes.

## **Letra Capitular**

Insere uma letra maiúscula grande no início do parágrafo, para dar-lhe, talvez, a aparência jornalística...

## ***Linha de Assinatura***

No mínimo inusitada, essa ferramenta cria automaticamente uma “linha” para o leitor poder assinar.

Mas, calma... Não é tão simples assim! Essa assinatura não será feita “no papel” e sim no próprio documento, por meio de assinatura digital (será necessário apresentar uma credencial digital). É um processo muito interessante: o autor cria o documento (vamos supor que o autor seja o proprietário de um imóvel e que o documento é o contrato de locação do mesmo); o leitor (que será o locatário, claro) lê o documento, confirma que está tudo OK e, usando seu certificado digital, assina o documento, salvando-o pela última vez.

Essa cópia do arquivo será devolvida ao autor para que este a guarde como um contrato em papel devidamente assinado e com firma reconhecida. Sim! Assinatura digital tornou-se muito mais simples no Office 2010!

## **Data e Hora**

Insere a data e a hora atuais. É possível escolher diversos formatos de data e hora como: sexta-feira, 11 de janeiro de 2008; 11/1/2008 16:00:59; 11/1/2008; jan-08;

## **Objeto**

Insere objetos no documento. Objeto é qualquer recurso (textos, fotos, imagens, músicas, sons, vídeos, planilhas, slides etc.) que o Windows consegue entender. É possível colocar basicamente qualquer tipo de coisa no Word.

Há outra opção, também, chamada “texto do arquivo”, que permite inserir um arquivo que o Word compreenda dentro do documento atual, para que o texto (conteúdo) do arquivo que foi inserido passe a constar no documento atual (um “enxerto” que até poderia ser feito via “copiar e colar”).

Veja alguns dos itens mencionados aqui inseridos devidamente no Word...

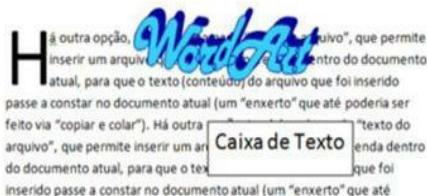


Figura 6.43 – Alguns recursos do Grupo Texto.

#### 6.2.2.7. Grupo Símbolos

##### Equação

Permite inserir e editar uma fórmula matemática com todos os recursos visuais possíveis (frações, raízes, expoentes, matrizes, índices, somatórios, derivadas e integrais).

##### Símbolo

Inserir símbolos. Símbolos são caracteres especiais existentes nos tipos de fonte instalados no Windows e normalmente não acessíveis pelo teclado diretamente. Com essa ferramenta, um símbolo  $\beta$  ou  $\text{€}$  pode ser inserido em qualquer lugar onde o Ponto de Inserção estiver localizado.

#### 6.2.3. Guia Layout da Página

A guia Layout da Página reúne os comandos relacionados com as configurações do documento em relação à página, como margens, cor de fundo, entre outros...



Figura 6.44 – Guia Layout da Página.

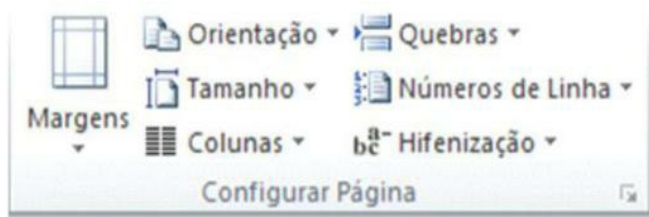
##### 6.2.3.1. Grupo Temas

O grupo Temas reúne comandos relacionados ao ajuste de Temas no documento. Um tema é um conjunto de definições de formatos para textos e formas no documento. A partir da utilização de um tema, pode-se mudar completamente a “cara” do documento, incluindo tipos de letras, cores, efeitos de formas (como 3D e sombra), em apenas alguns cliques.

Há vários temas predefinidos no Word 2010 e o usuário ainda poderá criar seus próprios temas!

#### 6.2.3.2. Grupo Configurar Página

Este grupo reúne ferramentas relacionadas com o formato da página em si. É basicamente o mesmo que consegue o comando Arquivo/Configurar Página das versões anteriores do Word (2003 e mais antigos). Inclusive, a setinha na parte inferior deste grupo abre a própria janela para Configurar Página.



*Figura 6.45 – Grupo Configurar Página.*

#### Margens

Configura as margens da página de acordo com alguns modelos predeterminados, como Estreita, Espelhada, Moderada, entre outras... Além disso, permite-se que o usuário configure manualmente tais margens (opção Margens Personalizadas...).

#### Orientação

Permite escolher entre escrever na página em pé (retrato) ou deitada (paisagem).

#### Tamanho

Configura o tamanho (dimensões) da página (A4, A3, Carta, Ofício etc.). É possível definir tamanhos personalizados (diferentes dos tamanhos predeterminados pelo “mercado”).

#### Colunas

Apresenta opções para formatar o texto em colunas (essa opção era encontrada, nas versões anteriores do Word, no menu Formatar/Colunas).



## Quebras

Insere quebras manuais (de página, de coluna ou de seção) no texto.

Uma quebra é uma interrupção forçada no fluxo do texto, fazendo “saltar” o cursor para o início do próximo objeto quebrado (próxima página, por exemplo, numa quebra de página). Esse recurso era encontrado, nas versões anteriores, no menu Inserir.

Uma “quebra” muito comum é conseguida ao pressionar ENTER: a quebra de parágrafo (afinal, quando pressionamos ENTER, o ponto de inserção salta de onde está para o início do próximo parágrafo).

Outras quebras possuem teclas de atalho também:

- Quebra de página: **CTRL + ENTER**
- Quebra de coluna: **CTRL + SHIFT + ENTER**
- Quebra de linha: **SHIFT + ENTER**

## Números de Linha

Permite exibir ou ocultar números que identificam as linhas do documento. É possível deixar que os números sejam contínuos (contados desde o início do documento) ou que se reiniciem página a página (ou seção a seção).

Atenção! Esse recurso fica visível apenas NA TELA do computador, enquanto se edita o documento (ou seja, os números de linhas não são impressos).

## Hifenização

É um recurso útil, mas proble-mático, de separação de sílabas. Normalmente não o utilizamos porque, nas versões anteriores do pro-grama Word, ele não era instalado por padrão. Esse recurso permite que, em vez de jogar uma palavra inteira para a próxima linha quando ela não cabe (esse é o comportamento do Word, como você bem sabe), o Word faça a separação silábica (como você pode perceber neste parágrafo!).

A hifenização (ops, hifenização) pode ser automática (ligada o tempo todo, no documento todo, enquanto se digita) ou manual (aplicada ao texto selecionado). Na hifenização manual, o Word notará que uma palavra pode ser hifenizada e sugerirá qual a melhor disposição para ela no texto (ou seja, em qual sílaba ocorrerá a colocação do hifen). O recurso de hifenização é usado para o layout (organização visual) dos textos jornalísticos.

### 6.2.3.3. Grupo Plano de Fundo da Página

Este grupo de comandos permite alterar as configurações visuais do fundo da página, como o:

#### Marca D'água

Insere um texto (ou figura) bem transparente atrás do texto do documento (em todas as páginas do documento);

#### Cor da Página

Parece absurdo que essa opção exista, especialmente se pensarmos em “imprimir” o documento! Mas é possível, sim, escolher uma outra cor para o “papel” no Word (esse recurso é particularmente útil quando se utiliza o Word para criar páginas da Web!).

## Bordas da Página

Esse recurso, que nas versões anteriores era encontrado em Formatar/Bordas e Sombreamento, serve para criar uma borda ao redor das páginas do documento.

Só um detalhe, caro leitor: se você está achando que é muita coisa para “decorar”, então, NÃO DECORE! Vá mexer no Word! Acompanhado deste livro, vá a cada opção aqui mostrada e veja como elas são!

Vai ficar muito mais fácil de lembrar e entender os itens!

E não se esqueça do curso em vídeo de Word 2010 no Eu Vou Passar ([www.euvoupassar.com.br](http://www.euvoupassar.com.br)). Ele traz detalhadamente cada um desses itens!

### 6.2.3.4. Grupo Parágrafo

Este grupo de ferramentas traz recursos encontrados na antiga janela de Formatar/Parágrafo – preste atenção a isto: a setinha que se apresenta neste grupo abre a mesma janela que a setinha do grupo *Parágrafo* da guia *Página Inicial*!

## Recuar

Controle de recuo. Esse recurso permite definir os recuos à esquerda e à direita do texto (novamente: recuo é o afastamento entre o texto e a margem do documento).

## Espaçamento

Determina os espaços que serão dados antes e depois do parágrafo (não é espaçamento entre linhas! É espaçamento antes do parágrafo e depois dele).

### 6.2.3.5. Grupo Organizar

Neste grupo são colocadas as ferramentas que trabalham com os objetos do documento (são considerados objetos fotos, organogramas, gráficos, autoformas, WordArt entre outros). Portanto, as ferramentas deste grupo só ficam disponíveis (habilitadas para uso) se algum objeto estiver selecionado (se um texto estiver selecionado, não há habilitação destas ferramentas).



Figura 6.46 – Grupo Organizar.

## **Posição**

Permite determinar em que posição, em relação à página, o objeto selecionado ficará (superior, inferior, meio, direita, esquerda etc.). É uma ferramenta muito simples de usar e em apenas um único clique, o objeto será posicionado no local devido.

## **Avançar**

Coloca o objeto selecionado à frente dos demais objetos da página (este comando só funciona com autoformas, ou seja, se outro tipo de objeto estiver selecionado, ele não habilitará esta ferramenta).

É possível escolher entre três opções: Avançar (que coloca a autoforma à frente do objeto imediatamente à frente dela atualmente), Trazer para a Frente (que coloca a autoforma à frente de todas as demais) e Trazer para a Frente do Texto (que coloca a autoforma à frente do texto do documento).

## **Recuar**

Coloca o objeto selecionado atrás de todos os demais objetos do documento.

Também há três opções neste comando: Recuar (que põe a autoforma atrás de quem atualmente está atrás dela imediatamente), Enviar para Trás (que a envia para trás de todas as demais autoformas) e Enviar para Trás do Texto (que coloca a autoforma atrás do texto do documento).

## **Quebra de Texto Automática**

Determina como uma figura (imagem, autoforma, gráfico, caixa de texto qualquer) vai se comportar em relação ao texto que a rodeia. É possível fazer com que o texto simplesmente circunde o objeto (escolhendo a opção Quadrado) ou fazer com que o objeto fique atrás do texto.

## **Alinhar**

Permite organizar os objetos entre si, alinhados por sua linha de base ou por sua linha superior. Ainda é possível alinhá-las pelos seus centros ou à esquerda/à direita. Note que não é o alinhamento do texto em relação à página, e sim dos objetos (fotos, imagens, autoformas) em relação aos demais objetos ou em relação à página.

## **Agrupar**

Reúne vários objetos em grupos, para que possam ser manipulados (mover, aumentar, diminuir etc.) sempre juntos.

Dentro deste comando, também, existem as opções *Desagrupar* e *Reagrupar*.

Este comando só está selecionável se houver mais de um objeto simultaneamente selecionado ou se um grupo estiver selecionado.

## **Girar**

Permite que os objetos selecionados sejam girados (rotacionados) ao redor de um ponto específico. Essa ferramenta permite, inclusive, determinar esse ponto fixo ao redor do qual o objeto vai girar.

#### 6.2.4. Guia Referências

Traz todos os recursos necessários aos dados que serão usados como referência no texto, como sumários, bibliografias, referências cruzadas e notas de rodapé.



*Figura 6.47 – Guia Referências.*

#### 6.2.4.1. Grupo Sumário

Este grupo apresenta recursos relativos à criação de sumários (índices de conteúdo).

É possível criar sumários automáticos desde que se tenha definida uma estrutura de estilos de tópicos (Título 1, Título 2, Título 3 etc.).

Depois de pronta a estrutura e devidamente aplicada aos capítulos do documento, basta acionar o botão Sumário para escolher e aplicar o estilo desejado de sumário.

O botão Atualizar Sumário permite que o Word reflita no sumário (modificando-o) qualquer alteração feita no documento (como mudança de páginas, e alteração nos textos dos títulos para os quais o sumário aponta).

#### 6.2.4.2. Grupo Notas de Rodapé

Este conjunto de ferramentas é usado para configurar e aplicar as notas de rodapé (aparecem na parte inferior da página onde a nota é inserida) ou notas de fim (aparecem no final do documento). A setinha no canto inferior direito abre uma janela de configuração de notas de rodapé e de fim.

É possível inserir uma nota de rodapé na página em questão através do botão **Inserir Nota de Rodapé**. Essa nota é um pequeno comentário que aparecerá no final da página onde a indicação da nota estará presente. A nota é indicada, no texto, por meio de um pequeno número, como o mostrado na figura a seguir).

Testando a Nota de rodapé<sup>1</sup>.

Vamos colocar<sup>2</sup> a segunda nota.

---

<sup>1</sup> Nota de Rodapé cria um índice (sobrescrito) na posição do Ponto de Inserção no texto. Além disso, coloca, no final da página, o número deste índice e a explicação referente a ele.

<sup>2</sup> Colocam-se as notas em Referências / Notas de Rodapé

Figura 6.48 – Notas de Rodapé.

O botão **Inserir Nota de Fim** serve para, digamos, inserir uma nota de fim!(É nada! Mesmo!?! Nem desconfiava, hein?) – uma nota de fim é um comentário que será colocado, pelo Word, **no final do documento** inteiro.

O botão **Próxima Nota de Rodapé** faz o Word remeter-se para a próxima (ou anterior) nota existente no texto (apesar do nome, ele faz a busca entre as notas de rodapé e as de fim).

O botão **Mostrar Notas** faz o Word exibir as áreas onde as notas de rodapé e de fim estão localizadas.

#### 6.2.4.3. Grupo Citações e Bibliografia

Eita negócio interessante (para quem cria apostilas, livros e/ou trabalhos de faculdade e dissertações)!

Com as ferramentas presentes neste grupo, é possível gerenciar as fontes bibliográficas usadas para o seu texto, bem como inseri-las da maneira adequada no texto (seguindo vários padrões mundiais de bibliografia).

Em primeiro lugar, usa-se o botão **Gerenciar Fontes Bibliográficas** para determinar, num banco de dados interno ao documento, quais são as fontes de bibliografia consultadas para a redação do documento. Depois, é possível determinar citações (no botão **Inserir Citações**) e inseri-las em pontos específicos do documento.

“Este é um exemplo de texto que pode ter sido retirado de algum livro... E o pequeno texto que

vem depois do fechamento das aspas é uma citação automática.” (João Antonio, 2012)

Através do botão **Bibliografia**, será construída, de forma automática (como o sumário), uma bibliografia no final do documento. Esse recurso só poderá ser usado se houver fontes bibliográficas devidamente inseridas e configuradas pelo botão Gerenciar Fontes Bibliográficas. Eis um exemplo de registro da bibliografia criada pelo Word:

Antonio, João (2006). *Informática para concursos* – 3ª edição. Rio de Janeiro – RJ: Campus/Elsevier.

É massa, não?!

#### 6.2.4.4. Grupo Legendas

Este grupo reúne ferramentas relacionadas à gerência das legendas das figuras e tabelas (pequenos textos que acompanham as figuras e tabelas, normalmente abaixo delas).

Ainda é possível, por meio do botão **Inserir Índice de Ilustrações**, criar um índice para as figuras no documento (figuras que possuem legendas, claro).

O botão **Referência Cruzada** permite inserir, num determinado trecho do texto, um link (vínculo) com outra parte do documento, como “como visto na figura 10.1”... É automático e muito simples. Caso a posição da figura mude no decorrer do texto, essa referência será automaticamente atualizada.

#### 6.2.4.5. Grupo Índice

Permite inserir índices remissivos (índices que apontam para as palavras-chave que aparecem no texto). Esse grupo **não permite** a criação de índices analíticos (sumários) nem índices de ilustrações (figuras).

#### 6.2.5. Guia Correspondências

Todas as ferramentas desta guia são relacionadas da mala-direta, etiquetas de endereçamento e preenchimento de envelopes de correspondência.



Figura 6.49 – Guia Correspondências.

##### 6.2.5.1. Grupo Criar

#### Envelopes

Permite a criação de envelopes (o Word editará um documento com o tamanho exato de um envelope) – é necessário, porém, que a impressora tenha suporte a envelopes para que o negócio funcione.

**Etiquetas**

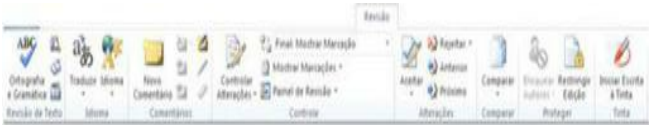
Permite a criação de um documento de etiquetas (ideal para trabalhar com aquelas etiquetas autoadesivas vendidas em folhas adesivas especiais).

**6.2.5.2. Demais Grupos**

As demais ferramentas referem-se ao recurso de Mala Direta, através do qual se pode preencher uma série de documentos quase idênticos, que variam apenas em alguns pontos (como uma mesma carta que será enviada a várias pessoas, modificando, de uma para a outra, apenas o nome do destinatário, por exemplo);

**6.2.6. Guia Revisão**

Como o nome já diz, essa guia reúne, em grupos, os comandos relacionados com o processo de revisão e correção do texto, como a correção ortográfica e gramatical, a tradução (sim... Tradução!), os dicionários de sinônimos, os comentários, o controle de alterações, entre outros...



*Figura 6.50 – Guia Revisão.*



*Figura 6.51 – Detalhe dos Grupos.*

**6.2.6.1. Grupo Revisão do Texto**

As ferramentas principais de revisão do texto encontram-se aqui.

## **Ortografia e Gramática**

Mesmo recurso amplamente conhecido de versões anteriores do Word. Vasculha o texto à procura de erros gramaticais (marcados em verde) e ortográficos (em vermelho). Assim como nas versões anteriores, esse comando é acessado pela tecla de atalho **F7**.

## **Pesquisar**

Abre o painel Pesquisar (que aparece à direita da página do documento) para que se possam encontrar materiais acerca do trecho desejado. O trecho é pesquisado em vários materiais de referência, como dicionários, livros on-line, serviços de tradução na Internet etc.

## **Dicionário de Sinônimos**

Sugere sinônimos para as palavras selecionadas.

## **Contar Palavras**

Recurso para apresentar a contagem de caracteres, palavras, parágrafos e outros itens do texto. (lembre-se: não conta apenas palavras, conta vários itens no texto)

### **6.2.6.2. Grupo Idioma**

#### **Traduzir**

Traduz o trecho selecionado para outro idioma (não confie! Nem sempre funciona 100%).

Ao clicar nesta opção, é possível traduzir o trecho selecionado, traduzir o documento inteiro ou ligar o ***Minitradutor***, que apresentará uma dica rápida de tradução para o idioma selecionado sempre que o mouse estiver em alguma palavra.

#### **Idioma**

Define, para um trecho de texto, qual é o idioma que se está usando para que o Word proceda à forma adequada de correção ortográfica e gramatical.

### **6.2.6.3. Grupo Comentários**

Este grupo trabalha com (como direi?) comentários! Comentários são lembretes que o usuário pode anexar a um documento (estes lembretes ficarão escritos em “quadros” lateralmente dispostos na página).

Para que serve um comentário? Simples! Para que um autor possa dar “recados”, por exemplo, para um coautor de um livro... Coisas do tipo “isso aqui é assim mesmo?” ou “lembre-se de pesquisar a bibliografia”.



Testando a Nota de rodapé<sup>1</sup>.

Vamos colocar<sup>2</sup> a **segunda** nota.

[30] Comentário tem certeza de que é o primeiro comentário mesmo?

<sup>1</sup> Nota de Rodapé cria um índice (subescrito) na posição do Ponto de Inserção no texto. Além disso, coloca, no final da página, o número deste índice e a explicação referente a ele.

<sup>2</sup> Colocam-se as notas em Referências / Notas de Rodapé

I

Figura 6.52 – Comentário vinculado à palavra “segunda”.

Se vários comentários de usuários diferentes estiverem no mesmo documento, eles serão apresentados em cores diferentes para cada usuário.

No Word 2010, é possível até mesmo inserir comentários “escritos” (além dos já conhecidos comentários “digitados”), que faz uso dos recursos dos tablets (caneta à mão).

#### 6.2.6.4. Grupo Controle e Grupo Alterações

As ferramentas contidas nestes dois grupos tratam das configurações e da utilização do recurso conhecido como **Controle de Alterações**.



Figura 6.53 – Grupos Controle e Alterações.

Através dos recursos presentes aqui, o usuário pode configurar para, por exemplo, quando alguma alteração for feita no documento, ela ficar demonstrada no próprio documento, como textos adicionados serem colocados em vermelho e sublinhados, ou textos excluídos serem

colocados em azul e tachados.

As ferramentas do grupo Alterações permitem navegar entre as alterações do texto e aceitá-las ou rejeitá-las.

Através dos botões **Aceitar** e **Rejeitar**, é possível acatar (ou não) uma determinada alteração no trecho selecionado ou até mesmo todas as alterações do documento (clitando no botão adequado, abre-se um menu com a opção de “aceitar todas” ou “rejeitar todas”).

#### 6.2.6.5. Demais Grupos

Eis os grupos e ferramentas faltantes na guia Revisão:



Figura 6.54 – Comparar, Proteger e Tinta.

#### Comparar

Esta ferramenta (única no grupo dela) oferece a capacidade de comparar dois documentos do Word, gerando, ao final do processo, um terceiro apresentando as diferenças entre eles como meras alterações.

Dentro dela, ainda é possível acionar o comando **Combinar**, que funde dois documentos do Word, em um terceiro, contendo todo o conteúdo dos dois documentos mesclados num único.

Também lá dentro, o comando **Mostrar Documentos de Origem** serve para exibir quais foram os dois documentos que deram origem àquele documento comparado (ou combinado) que se está visualizando naquele momento.

#### Proteger

Os comandos neste grupo permitem limitar o poder de manipulação do arquivo, restringindo privilégios e autores (ou seja, determinando “quem pode” e “o que pode” fazer no documento).

O botão **Restringir Edição**, ao ser clicado, abre um painel que aparece à direita do documento. Neste painel, as opções de restrição e proteção se apresentam.

É possível determinar áreas editáveis e bloquear todo o restante do documento, por exemplo. É possível, também, impedir que se escreva diretamente dentro do documento, permitindo só a alteração dos comentários.

## Tinta

O botão **Iniciar Escrita à Tinta** permite que o usuário escreva (com a caneta do tablet ou mesmo com o mouse) diretamente no corpo do documento.

Diferentemente do “Comentário à Tinta” (que permite a escrita via caneta dentro de um retângulo daqueles do comentário), esta ferramenta permite a escrita diretamente na página (no meio do texto mesmo).

Há como escolher entre Caneta, Marca-Texto e Borracha (para, claro, apagar as anotações feitas com a caneta e o marca-texto).

### 6.2.7. Guia Exibição

Essa guia conta com comandos que eram (nas versões anteriores do Word), em sua maioria, localizados no menu **Exibir**, como é o caso do Zoom e dos Modos de Exibição do documento. Há também alguns comandos de outros locais, como Macro (antigamente no menu Ferramentas) e recursos usados em Janelas.



Figura 6.55 – Guia Exibição.

#### 6.2.7.1. Grupo Modos de Exibição de Documento

Esse grupo permite alternar entre os possíveis modos de exibição do documento que está sendo editado. Esses modos de exibição apenas modificam a forma como o documento é visto, e não seu conteúdo em si.

As opções de exibição de documento são as seguintes:

##### Layout de Impressão

É o formato mais usado, em que vemos a página em branco na nossa frente (inclusive vemos as margens do papel, o cabeçalho e o rodapé).

Este é o modo que oferece WYSIWYG (What You See Is What You Get – ou seja, “O que você está vendo é o que você vai obter”) – em suma, neste modo de exibição, você consegue trabalhar visualizando EXATAMENTE (ou o mais aproximado possível) aquilo que vai sair impresso!

Atenção: não é possível ver o cabeçalho e o rodapé em outro modo de exibição! Inclusive, é bom que se saiba: quando estamos em outro modo de exibição qualquer e o usuário pede para ver (ou editar) o cabeçalho ou o rodapé do documento, este passa a ser exibido, momentaneamente, no modo Layout de Impressão!

## Leitura em Tela Inteira

É o formato mais indicado para quem quer ler o documento na tela do computador. O documento é apresentado em telas divididas como um livro no monitor. Dá até para “passar” as páginas como se faz num livro normal.

## Layout da Web

Melhor maneira de visualizar o documento, se ele estiver sendo desenvolvido para a Internet (ou seja, se você estiver fazendo uma página da Web com o Word).

## Estrutura de Tópicos

Este modo permite que se visualize o documento resumido, dividido e organizado por sua estrutura de tópicos (títulos) de modo que se possa facilmente passar de um nível para outro sem ter que passar por uma quantidade absurda de texto.

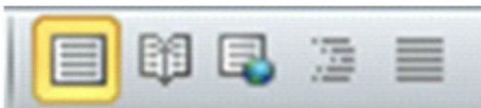
Você pode “contrair” ou “expandir” tópicos para visualizar e navegar mais facilmente no texto.

## Rascunho

Antiga exibição Normal (nas versões anteriores do Word). Esse modo de exibição mostra o documento sem margens ou páginas. A tela inteira fica branca e a divisão entre as páginas é mostrada como uma linha tracejada.

As figuras do documento, por padrão, não são mostradas (isso não acontecia nas versões anteriores). A justificativa para isso é que o modo Rascunho permite uma edição “rápida” do documento, sem que se precise carregar na memória os itens mais pesados (e, claro, que deixarão o micro mais lento).

Também é possível alterar os modos de exibição do Word clicando-se nos botões adequados na **barra de status** do Word (na parte inferior do programa) – na extremidade direita da barra de status. Em ordem, são eles: Layout de Impressão, Leitura em Tela Cheia, Layout da Web, Tópicos e Rascunho (ou seja, eles têm a mesma organização dos botões do Grupo).



*Figura 6.56 – Botões dos Modos de Exibição na barra de status.*

### 6.2.7.2. Grupo Mostrar

Esse grupo, composto apenas de checkboxes, ou caixas de verificação (é esse o nome dado a esses “quadrinhos” onde podemos clicar para aparecer o sinal de “visto” ou “check”), permite definir o que será mostrado na interface do Word e o que não será.

Vários componentes podem ser exibidos ou ocultos (quando o “quadrado” está marcado, é

sinal de que aquele recurso vai ser mostrado... Acho que você já havia entendido isso, não é mesmo, caro leitor?!).

## Régua

Permite mostrar/ocultar as régua horizontal e vertical do documento (aquelas que aparecem acima e à esquerda do documento, respectivamente);

Se o seu documento já está mostrando as régua, é porque este checkbox está marcado!

## Linhas de Grade

Diz respeito às linhas que delimitam a estrutura de uma tabela. Se você não aplicar nenhuma formatação de bordas numa tabela, mesmo assim, ela aparecerá no documento da tela, em linhas cinza muito tênues... São as linhas de grade! Se você escolher não exibi-las, a tabela inteira parecerá invisível (só o conteúdo será visto na tela, não as linhas em si).

Note bem: as linhas de grade das tabelas não serão impressas! Mesmo que você peça para exibi-las na tela, elas não serão colocadas no papel! Para que uma tabela seja vista no documento impresso, é necessário formatar suas bordas, pintando-as.

## Painel de Navegação

Exibe/oculta o painel de navegação, na lateral esquerda da janela do Word. Com esse painel aberto, é possível navegar (ir e vir) pelo documento facilmente usando a sua estrutura de títulos (tópicos) – só funciona bem se os títulos dos capítulos tiverem sido formatados com estilos hierárquicos (título 1, título 2 etc.).



Figura 6.57 – Painel de Navegação aberto.

### 6.2.7.3. Grupo Zoom

Esse é um dos mais fáceis, pois diz respeito apenas ao zoom, que é a aproximação do documento em relação ao usuário. O zoom pode aproximar (aumentando seu valor) ou afastar (diminuindo seu valor) o documento do usuário.

É possível, também, escolher entre valores que permitam a visualização de mais de uma página simultaneamente.

O zoom também pode ser conseguido em sua barra deslizante, na extremidade inferior direita da janela do Word (na barra de status, logo depois dos botões de modo de exibição).



*Figura 6.58 - Barra deslizante do Zoom na barra de status do Word.*

#### **6.2.7.4. Grupo Janela**

Esse grupo traz uma série de ferramentas relacionadas ao trato com várias janelas simultaneamente. Os recursos deste grupo servem para conseguir visualizar duas, ou mais, partes de um mesmo documento.

Desta forma, é possível digitar em um trecho específico consultando (lendo) o que está escrito em outra parte do documento.

##### **Nova Janela**

Abre uma janela diferente, mas do mesmo documento. Ou seja, será possível ver duas janelas de um mesmo conteúdo (mesmo arquivo).

##### **Dividir**

Cria uma linha divisória para separar a janela em dois painéis (um em cima do outro). Os painéis são independentes e permitem visualizar simultaneamente duas partes de um mesmo documento.

Esse recurso serve para o mesmo objetivo do comando Nova Janela, mas ele não abre uma nova janela: os dois painéis criados ficam dentro de uma única janela.

##### **Organizar Tudo**

Organiza a exibição de todas as janelas do Word abertas (mesmo que elas sejam de documentos diferentes). As janelas serão redimensionadas (as que estiverem maximizadas são restauradas para baixo) de modo que na tela do computador se possa ver todas elas!

As janelas abertas ficarão organizadas uma em cima da outra!

##### **Exibir Lado a Lado**

Organiza duas janelas de forma que cada uma ocupe metade da tela. As janelas ficam organizadas uma do lado da outra (dããã! Deu pra deduzir, né?).

Quando este recurso está ativado, automaticamente, habilitam-se os comandos **Rolagem Sincronizada** (que faz com que o uso da barra de rolagem de uma das janelas influencie na rolagem das duas ao mesmo tempo) e **Redefinir Posição da Janela** (para fazer as janelas ocuparem, novamente, 50% da tela cada uma – caso você tenha redimensionado alguma delas).

#### 6.2.7.5. Grupo Macros

Permite a manipulação de macros. Macro é o nome dado a um pequeno programa (conjunto de comandos listados em ordem sequencial) usado para automatizar tarefas no Word. Uma macro pode fazer desde formatação automática (como aplicação de vários efeitos simultaneamente) até a construção inteira de documentos, tabelas, índices etc.

Na boa, caro leitor, acho que este comando só veio parar aqui porque o pessoal lá da Microsoft não sabia exatamente onde colocá-lo! Do tipo: “Terminamos de distribuir as ferramentas entre as guias!” – “Mas, e MACRO? Tá faltando!” – “Ihhh, cara! Coloca em qualquer canto! Enfia aí na guia Exibição mesmo, senão a gente vai ter que fazer tudo de novo!”.

#### 6.2.8. Ferramentas e Guias Interativas

Apesar de termos visto basicamente todas as guias e grupos presentes na faixa de opções do Word 2010, é bom que você saiba, caro aluno, que há mais guias, que só aparecem em determinados casos: essas são as Guias Interativas (elas só aparecem se o contexto selecionado assim necessitar).

Essas guias aparecem na Faixa de Opções (na extremidade direita) e só ficam ativas e visíveis enquanto o objeto certo (aquele a que ela está relacionada) estiver selecionado.

##### 6.2.8.1. Ferramentas de Imagem

Esta guia aparece quando uma imagem é selecionada (tome-se por imagem um arquivo fotográfico inserido no documento do Word). Entram no conceito de imagem os arquivos com extensões JPG, GIF, TIFF, BMP, PNG, EMF, WMF e mais alguns...

Essa “guia”, na verdade, não é apresentada no mesmo nível das outras guias da faixa de opções. O título “Ferramentas de Imagem” aparece na barra de título e, abaixo dele, vem uma única guia chamada formatar.



*Figura 6.59 – Guia Formatar, das Ferramentas de Imagem.*

Nesta guia, é possível realizar operações com figuras (imagens importadas para o Word), como determinar brilho, contraste, bordas especiais (como molduras de foto, por exemplo), sombra, efeito de reflexo, cortar a imagem, determinar a posição em relação ao texto, entre outras coisas.

##### 6.2.8.2. Ferramentas de Desenho

Esta guia só se torna disponível quando um desenho estiver selecionado (tome-se por desenho uma autoforma, ou qualquer figura vetorial que possa ser entendida e editada como uma autoforma);



*Figura 6.60 – Guia Formatar, das Ferramentas de Desenho.*

Nesta guia, podemos inserir novas autoformas, alterar sombra, efeitos 3D, organizar (enviar para trás, enviar para frente), girar, agrupar e desagrupar autoformas e objetos afins.

### 6.2.8.3. Ferramentas de Tabela

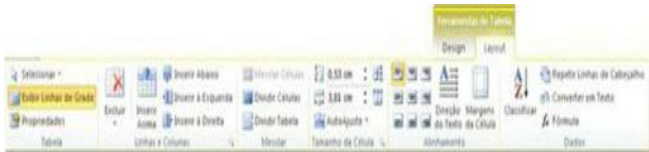
Precisa dizer alguma coisa? Este conjunto de ferramentas contém guias para manipulação de tabelas e, claro, só aparece quando alguma tabela estiver selecionada (ou quando o cursor estiver localizado em alguma célula da tabela). Há duas guias em Ferramentas de Tabela: Design e Layout.



*Figura 6.61 – Guia Design, das Ferramentas de Tabela.*

A guia Design (mostrada acima) é responsável pelos efeitos visuais da tabela (cores, linhas de borda etc.) – caso queira “enfeitar” a tabela, é aqui!



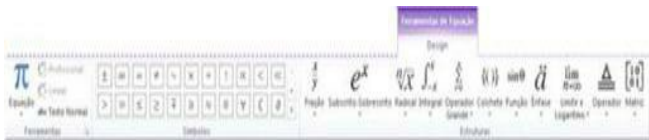


*Figura 6.62 – Guia Layout, das Ferramentas de Tabela.*

A guia Layout traz comandos referentes à estrutura da tabela, como opções para excluir linhas e/ou colunas, mesclar células, distribuir colunas ou linhas, classificar (ordem crescente ou decrescente), converter tabela em texto, aplicar fórmulas etc.

#### 6.2.8.4. Ferramentas de Equação

Como o nome já diz, esse conjunto de comandos apresenta recursos valiosíssimos para a edição de equações criadas pelo Word 2010! Esta guia aparece quando alguma equação é selecionada (ou seja, quando o usuário dá um clique em qualquer lugar numa equação presente no documento).



*Figura 6.63 – Guia Design, das Ferramentas de Equação.*

O mais importante é lembrar que, ao contrário das versões anteriores, no Word 2010 (na verdade, no 2007 também), as equações são editadas pelo próprio Word! Ou seja, não é necessário nenhum programa adicional para modificar as equações, como se fazia anteriormente.

Portanto, qualquer equação existente dentro do documento do Word será manipulada (editada) diretamente pelo Word, bastando, para isso, que o usuário dê apenas um único clique em qualquer parte da equação que deseja alterar.

#### 6.2.8.5. Ferramentas de Cabeçalho e Rodapé

Essa guia aparece quando o usuário aplica duplo clique no cabeçalho ou no rodapé (essa ação permite que o usuário edite-os). A guia Ferramentas de Cabeçalho e Rodapé apresenta uma série de opções que podem inserir data/hora, números de páginas e outros recursos na parte superior

(ou na inferior) das páginas do documento.



Figura 6.64 – Ferramentas de Cabeçalho e Rodapé – Guia Design.

Para voltar à área normal de texto, basta aplicar duplo clique em qualquer lugar na área de texto do documento (ou seja, fora das áreas do cabeçalho e do rodapé) ou simplesmente clicar no comando **Fechar Cabeçalho e Rodapé**, na extremidade direita da guia.

### 6.3. Guia Arquivo

Na extremidade esquerda da Faixa de Opções encontra-se a guia **Arquivo**. Dentro dela, há muitas opções que merecem nossa atenção!

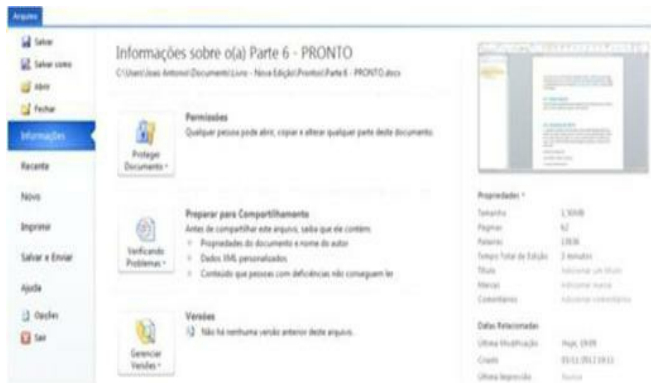


Figura 6.65 – Página de Informações da guia Arquivo.

Perceba que as opções contidas na tela da guia Arquivo estão localizadas na coluna à esquerda da tela. Vamos a esses itens agora:

### 6.3.1. Comandos de Arquivo

Logo na parte superior, podemos encontrar os comandos que podemos acionar para trabalhar com o arquivo. São eles:

#### 6.3.1.1. Salvar

Este comando, mais do que manjado, serve para salvar o documento que se está utilizando no momento.

Caso seja o primeiro salvamento, o Word lhe pedirá o nome do arquivo e o local onde ele será salvo (semelhante ao Salvar Como). Caso o arquivo já tenha sido salvo alguma vez, não lhe será pedida nenhuma informação.

A tecla de atalho para o comando Salvar é **CTRL + B**.

O Word 2010 salva seus arquivos, por padrão, no formato (extensão) DOCX. Mas o Word também consegue salvar arquivos em outros variados formatos, como DOC (versão antiga do Word), TXT (documento do bloco de notas) e até mesmo ODT (Texto Open Document, usado pelo BrOffice).

Ahhh! É bom que você saiba também: o Word 2010 consegue salvar arquivos diretamente no formato PDF (formato do Adobe Acrobat). É um grande avanço em relação às versões anteriores (o 2007 já fazia, mas os mais antigos, não!).

#### 6.3.1.2. Salvar Como

Este comando sempre pede o nome do arquivo a ser salvo e o local onde ele será salvo.

Não importa se é a primeira vez ou não: este comando vai sempre abrir uma janela, pedindo nome e local, e, com isso, ele permite que se criem novas cópias do documento em questão.

Se, ao utilizar o comando Salvar Como, você indicar o mesmo nome e local do arquivo atualmente, ele será salvo por cima do anterior (exatamente como o faz o comando Salvar).

A tecla de atalho para o comando Salvar Como é **F12**.

Claro que os formatos de arquivos usados são exatamente os mesmos do comando Salvar.

#### 6.3.1.3. Abrir

Este comando abre uma janela para escolher um arquivo previamente salvo para ser aberto na memória RAM do computador. Abrir é permitir que o conteúdo do arquivo seja visto e alterado pelo Word.

O comando Abrir pode ser acionado, também, pela tecla de atalho **CTRL + A**.

O Microsoft Word 2010 consegue abrir arquivos no formato DOC (extensão dos antigos formatos do Word), DOCX (formato atual), HTM e HTML (páginas da Web), RTF (Rich Text Format) e ODT (formato do Open Document, usado pelo BrOffice).

Portanto, não se esqueça: o Word consegue, sim, abrir arquivos do BrOffice (extensão ODT).

Ahhh! Um lembrete: o Word 2010 **não consegue abrir** arquivos no formato PDF (só consegue

salvá-los neste formato!) – para abrir arquivos PDF, deve-se usar o *Adobe Reader* (programa gratuito já visto).

#### 6.3.1.4. Fechar

Este comando fecha (retira da memória RAM) o arquivo que se está usando. O programa Word, em si, continua em execução, mas o arquivo que estava aberto passa a não estar mais!

Ao pedir para fechar um documento, se houve alterações recentes nele que não foram salvas, será perguntado se você deseja salvá-las agora!

A tecla de atalho para o comando Fechar é **CTRL + F4**.

#### 6.3.2. Páginas da guia Arquivo

Abaixo dos comandos, há uma série de itens que, na verdade, são atalhos para páginas (telas) específicas. Ao clicar em cada uma dessas opções, sua página própria aparece, assumindo o espaço da guia Arquivo. Vamos a elas:

##### 6.3.2.1. Informações

Por meio da página de Informações, é possível ter acesso a vários dados acerca do arquivo, como tamanho (em bytes), número de páginas, quantidade de palavras escritas, tempo total de edição etc.

Também é possível ter acesso a certos comandos como:

##### Proteger Documento

Oferece opções para: salvar o arquivo com senha de acesso e senha de proteção; criptografar o arquivo; adicionar assinatura digital e restringir edição.

##### Verificando Problemas

Permite analisar o documento para encontrar possíveis motivos de incompatibilidades com outros programas, caso seja necessário salvar o documento em formato diferente do Word.

##### Gerenciar Versões

Permite acesso a um recurso interessante que dá o direito de acessar e editar versões anteriores do mesmo documento (versões que existiam antes dos salvamentos), a fim, por exemplo, de encontrar trechos que já foram apagados em versões recentes.

##### 6.3.2.2. Recente

Apresenta uma lista com os documentos e locais recentemente manipulados (documentos, por exemplo, que você recentemente abriu e/ou salvou no seu computador).

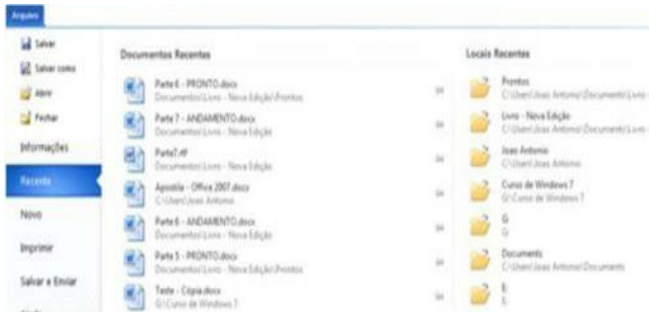


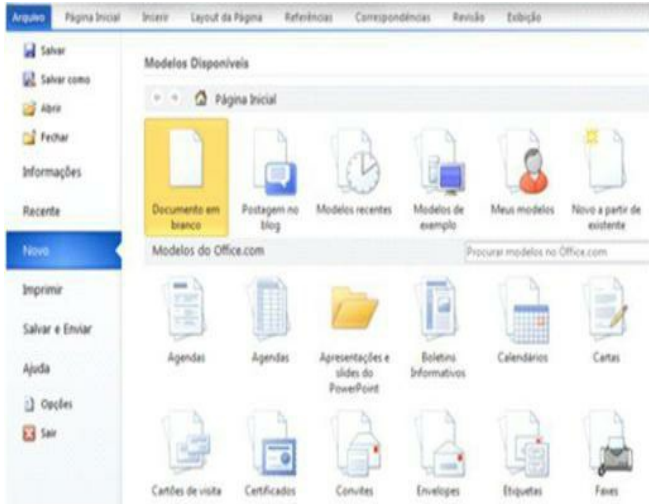
Figura 6.67 – (Sim, é a 6.67!) – Página recente da guia Arquivo.

### 6.3.2.3. Novo

Permite escolher tipos predefinidos de modelos para a criação de um novo documento na memória RAM.

Normalmente, quando criamos um novo arquivo, usamos a tecla de atalho **CTRL + O**, que automaticamente cria um documento novo em branco, baseado nas definições do modelo *Normal.dotx*.

Essa página, porém, nos dá, além do direito de criar documentos em branco (primeira opção da tela), direito de criar muitos outros tipos de documento, com base nos mais variados modelos.



*Figura 6.68 – Menu Novo, na guia Arquivo.*

#### 6.3.2.4. Imprimir

Esta página traz uma série de opções relacionadas, claro, ao ato de imprimir o documento. E é possível, claro, mandar imprimir o documento!

É possível escolher quantas cópias serão impressas; em qual impressora se vai imprimir; quantas e quais páginas vamos querer impressas; se vamos imprimir só as ímpares, só as pares, ou todas elas, entre outras opções.

Podemos configurar, também, definições relacionadas ao papel em si: margens, tamanho e orientação do papel, além do recurso de imprimir mais de uma página (do documento) por página (no papel).

No Word 2010, a tecla de atalho **CTRL + P** também abre a página Imprimir que estamos analisando.

#### 6.3.2.5. Salvar e Enviar

Oferece diversos recursos para salvar o documento do Word em formatos especiais (que também são encontrados no comando Salvar). Além disso, oferece recursos para já enviar o documento via e-mail ou via Web para outras pessoas ou provedores de conteúdo.

É possível, por exemplo, criar uma página da Web e já publicá-la num servidor apropriado na Internet por meio de opção contida aqui! Dá uma olhada!

#### 6.3.2.6. Ajuda

Oferece meios de responder as dúvidas do usuário na sua busca por aprender mais sobre o Word. Este menu também oferece recursos de acesso ao serviço de ajuda da Microsoft (online), por meio da Internet.

A tecla de atalho para abrir a Ajuda do Microsoft Office é **FI**.

### 6.3.3. Demais comandos

#### 6.3.3.1. Opções

Abre a janela de opções do programa Word. Esta janela é formada por muitas subseções. Qualquer coisa pode ser perguntada sobre esta janela, e seria uma covardia tremenda para com o candidato! ;-D

Eu sugiro que você nem se preocupe com essa janela, tá? Afinal, é muita coisa e, mesmo que você estudasse tudo, seria pouco provável que você conseguisse se lembrar do que estudou e é muito provável que caia exatamente aquilo que você não viu! ;-)

#### 6.3.3.2. Sair

Este comando fecha o programa Word.

Sua tecla de atalho é ALT + F4 (a mesma que fecha qualquer janela do Windows), afinal, esse comando é o equivalente a clicar no botão do X, lá na extremidade direita superior da janela!

### 6.3.4. Barra de Ferramentas de Acesso Rápido

Apesar de não pertencer, exatamente, à guia Arquivo, este conjunto de ferramentas merecia ser mencionado!

Já reparou um conjunto com alguns botões pequenos acima da guia Arquivo? Pois é, esta é a **Barra de Ferramentas de Acesso Rápido** e traz, por padrão, algumas ferramentas usadas quase que constantemente.



Na sequência, temos os comandos:

#### **6.3.4.1. Salvar**

Para salvar rapidamente o documento sem precisar acionar a guia Arquivo.

#### **6.3.4.2. Desfazer**

A cada clique nesta “setinha curva”, o Word desfaz uma ação realizada pelo usuário. Note a presença de uma setinha preta (apontando para baixo) à direita do botão desta ferramenta: ela serve para dar acesso à lista de ações que podem ser desfeitas no programa.

A tecla de atalho do comando desfazer é **CTRL + Z**.

#### **6.3.4.3. Refazer/Repetir**

O botão nesta foto (uma seta que dá um “giro” de 360 graus) é o botão repetir. Ele está aparecendo porque o comando Refazer não está habilitado.

O botão Refazer só é habilitado se alguma ação foi desfeita recentemente, ou seja, quando o botão desfazer acabou de ser usado.

O comando Refazer serve para “desfazer” o que o Desfazer realizou, ou seja, cancelar a desfeita. Já o comando Repetir serve para repetir o último comando realizado (por exemplo, um negrito, uma alteração de tamanho de fonte, etc.).

Ambos os comandos usam a mesma tecla de atalho: **CTRL + R**.

#### **6.3.4.4. Estilos Rápidos**

Este comando (na forma de uma “letra A com um pincel”), dá acesso à galeria de estilos rápidos (a mesma que aparece na guia Página Inicial, no grupo Estilo).

Um clique aqui já abrirá a galeria de estilos!

#### **6.3.4.5. Personalizar a Barra de Ferramentas de Acesso Rápido**

O último botão, com o formato de um “traço em cima de uma seta para baixo”, serve para abrir as opções de personalização da Barra de Ferramentas, permitindo adicionar e remover botões que você ache necessários!



#### 6.4. Questões de Word

1. (Analista/TJ-RS/2012) O MS-Word tem recursos de formatação de textos que podem ser usados simultaneamente, conforme o que se deseja destacar no texto. Porém, alguns deles não podem ser usados simultaneamente. Das alternativas abaixo, qual contém dois efeitos de formatação que podem ser usados simultaneamente?
  - a) Subscrito e Sobrescrito.
  - b) Versalete e Todas em maiúsculo.
  - c) Tachado e Tachado duplo.
  - d) Tachado e Subscrito.
  - e) Relevo e Baixo relevo.
2. (Técnico/TJ-PE/2007 – adaptada) No Word, quebras de páginas podem ser obtidas através da guia:
  - a) Página Inicial;
  - b) Inserir;
  - c) Layout da Página;
  - d) Referências;
  - e) Revisão.
3. (Auditor ISS-SP/2012) O MS Word:
  - a) é apenas um editor de textos, não permitindo a edição de figuras e tabelas.
  - b) não permite a construção automática de uma tabela de conteúdo para um documento.
  - c) possui recursos de correção ortográfica e correção gramatical.
  - d) permite a construção de slides com transições sofisticadas.
  - e) permite formatação condicional do documento, atribuindo-se fontes e cores de acordo com o seu conteúdo.
4. (MI/2012) No Microsoft Word:
  - a) pode-se copiar um texto através do recurso arrastar-e-soltar, mantendo-se a tecla Ctrl pressionada.
  - b) são efeitos de fonte: Tachado misto, Sobrescrito, Contorno, Relevância, Versalete.
  - c) pode-se copiar um texto através do recurso arrastar-e-soltar, mantendo-se a tecla Alt pressionada.
  - d) são efeitos de fonte: Tachado, Sobreposto, Compactado, Relevo, Versalete.
  - e) são efeitos de fonte: Tachado duplo, Inter-escrito, Contorno, Relevo, Versão.
5. (TJM/SP/2012) No MS-Word 2007, na guia Início, grupo Edição, existe o seguinte ícone.



Por meio desse ícone, denominado:

- a) colar, é possível copiar para o ponto de edição uma palavra ou frase da Área de Transferência.
- b) localizar, é possível procurar ocorrências de uma palavra ou frase específica.
- c) marcar, pode-se sublinhar palavras específicas no parágrafo selecionado.
- d) ordenar, pode-se colocar em ordem alfabética os elementos de uma tabela.
- e) selecionar, é possível marcar todo o texto do arquivo até a primeira ocorrência de uma determinada palavra.

6. (SEAP-RJ/2012) No Word 2010 BR, a finalidade do acionamento do ícone  e o atalho de teclado que corresponde ao ícone x<sup>2</sup> são, respectivamente:

- a) aumentar o nível de recuo de parágrafo e Ctrl + =
- b) aumentar o nível de recuo de parágrafo e Ctrl + +
- c) mudar o tamanho da fonte aplicada ao texto e Ctrl + +
- d) alterar o espaçamento entre as linhas de texto e Ctrl + +
- e) alterar o espaçamento entre as linhas de texto e Ctrl + =



7. (TCM-GO/2012 – adaptada) O ícone , extraído do Microsoft Word 2010, é associado a que finalidade?

- a) Iniciar uma mala direta.
- b) Enviar uma mala direta.
- c) Criar e imprimir envelopes.
- d) Selecionar destinatários para mala direta.
- e) Realçar os campos inseridos em um documento.

8. (Banco do Brasil/2012) No Microsoft Word versão 2007, para alinhar um texto selecionado tanto à margem direita quanto à margem esquerda, acrescentando espaço extra entre as palavras, conforme seja necessário, pode-se utilizar o atalho de teclado:

- a) Ctrl + E;
- b) Ctrl + J;
- c) Ctrl + D;

- d) Alt + A;
- e) Alt + B.

9. (MP-PE/2012) No Microsoft Word 2007 ou superior é possível salvar arquivos no formato de texto Open Document, usado por alguns aplicativos de processamento de texto, como o OpenOffice.org Writer e o Google Docs. A extensão de um arquivo salvo no formato de documento citado acima é:

- a) .odt;
- b) .pdf;
- c) .xps;
- d) .mdb;
- e) .pps.

10. (Banco do Brasil/2012) No Microsoft Word versão 2010, a guia Referências oferece, por padrão, o comando:

- a) visualizar Resultados, cuja finalidade é substituir os campos de mesclagem do documento pelos dados reais da lista de destinatários.
- b) definir Idioma, cuja finalidade é definir o idioma usado para verificar a ortografia e a gramática do texto selecionado.
- c) linhas de Grade, cuja finalidade é ativar linhas de grade para servir como referência no alinhamento dos objetos do documento.
- d) controlar Alterações, cuja finalidade é controlar as alterações feitas no documento, incluindo inserções, exclusões e alterações de formatação.
- e) marcar Citação, cuja finalidade é adicionar o texto selecionado como uma entrada no índice de autoridades.

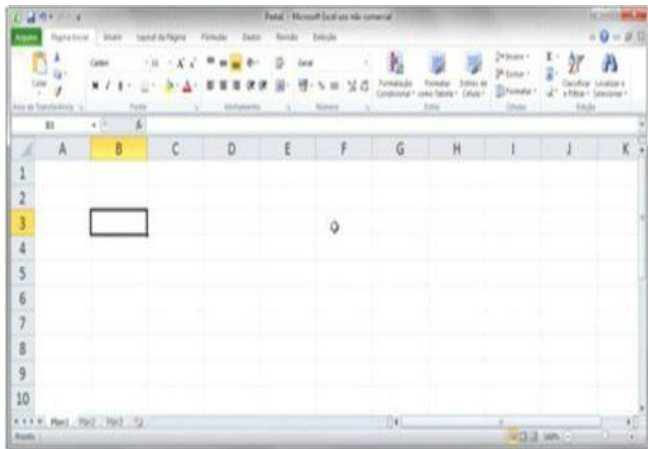
## Capítulo 7

### Microsoft Excel

#### 7.1. Conhecendo o Microsoft Excel

O Microsoft Excel é um programa gerenciador de planilhas eletrônicas de cálculos. Com ele, é possível criar tabelas numéricas para os mais diversos fins, desde simples calendários escolares a orçamentos completos de projetos dos mais variados tipos.

O Microsoft Excel é desenvolvido e comercializado pela Microsoft no pacote Office, e é acompanhado pelo Word e outros aplicativos diversos.



*Figura 7.1 – Janela do Microsoft Excel.*

A maioria das operações que serão vistas neste livro versa sobre a versão 2010 do Excel (a mais recente versão e, claro, a mais cobrada em prova atualmente). Mas não se preocupe, a maioria dos comandos e recursos cobrados no Excel é idêntica em todas as versões (cálculos e funções são os assuntos mais cobrados).

#### 7.2. Interface do Excel

A parte superior da janela do Excel é muito semelhante ao Word (a faixa de opções também é apresentada no Excel). Os comandos que o Excel possui, porém, são um tanto diferentes (claro, o Excel é outro programa, né?).

### 7.2.1. Faixa de Opções

Contém os comandos do Excel dispostos em guias separadas por conjuntos que conhecemos como Grupos. Algumas destas guias possuem grupos e ferramentas semelhantes ao Word.



*Figura 7.2 – Faixa de Opções do Excel (mostrando a guia Página Inicial).*

Para acessar as guias e ferramentas da Faixa de Opções, pressione a tecla ALT (isso funciona no Word também!). Você verá cada comando e guia sendo associado a uma letra que, se pressionada, irá acionar o respectivo comando.



*Figura 7.3 – Faixa de Opções depois da tecla ALT.*

### 7.2.2. Barra de Status

É a barra horizontal, localizada na base da tela do Excel que apresenta várias informações a respeito do estado da janela do programa. Consultar a barra de status do Excel é muito importante, embora não tanto quanto no Word.



*Figura 7.4 – Barra de Status do Excel.*

“Ei João, o que são aqueles números? Média, Contagem, Soma? Nunca notei aquilo.”

É muito simples, caro leitor. Quando há várias células selecionadas, a média, a contagem e a soma das mesmas aparece na barra de status, como mostrado na figura anterior.

É automático! Basta selecionar as células, não precisa fazer mais nada!

Se você clicar com o botão direito do mouse em qualquer parte da barra de status do Excel 2010, um submenu aparece, permitindo personalizar a barra de status e, entre suas opções, está a de escolher que tipo de informações destes autocálculos vão aparecer! Ahh! Só para constar, é claro que esse recurso é chamado de *Autocálculo*.

### 7.2.3. Caixa de nome

É uma área localizada acima da planilha que mostra o endereço da célula atual.



Figura 7.5 – Caixa de nome mostrando a célula C2.

A caixa de nome também permite ir diretamente para uma célula qualquer da planilha ou atribuir um nome a uma célula.

“Atribuir um nome, João? Como assim?”

Muito simples, caro leitor: é possível atribuir um nome amigável (fácil de lembrar) para qualquer célula da planilha, desde que não haja outra célula, no arquivo, com o mesmo nome. Para atribuir um nome a uma célula, basta selecionar a célula e digitar o nome na Caixa de Nome.

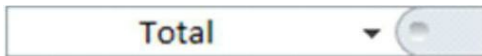


Figura 7.6 – Caixa de nome mostrando o nome atribuído a uma célula da planilha.

### 7.2.4. Barra de fórmulas

É a barra branca grande localizada acima da planilha que mostra o real conteúdo da célula selecionada. Caso o conteúdo da célula seja um cálculo (cujo resultado aparece na planilha), será apresentado nesta área o cálculo, não o resultado.

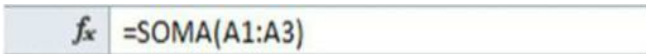


Figura 7.7 – A barra de fórmulas SEMPRE mostra o conteúdo real da célula selecionada.

Também é possível editar uma célula inserindo o conteúdo diretamente na barra de fórmulas. Para isso basta selecionar a célula que se deseja modificar, clicar uma única vez na barra de

fórmulas, digitar o que se deseja e confirmar a alteração (com ENTER).

### 7.2.5. Guias das planilhas

É uma área localizada na parte inferior da planilha que mostra a planilha atual de trabalho (ou seja, a planilha que estamos usando agora). Quando criamos um documento novo no Excel, ele normalmente fornece três planilhas independentes para trabalharmos.

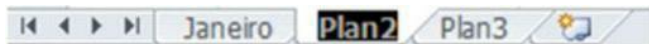


*Figura 7.8 – Guias das planilhas.*

A quarta guia mostrada (depois da Plan3, com o ícone estranho) serve para inserir uma nova planilha! Basta clicar nele e o Excel irá criar uma planilha após a Plan3.

Também é possível inserir novas planilhas por meio do comando **Inserir**, que fica no grupo **Células**, da guia **Página Inicial**.

É possível renomear uma planilha aplicando um duplo clique na guia desejada e digitando o novo nome. Para renomear uma planilha pela Faixa de Opções, acione o comando **Formatar**, também localizado no grupo Células, na guia Página Inicial.



*Figura 7.9 – Planilha Janeiro e planilha Plan2 sendo renomeada.*

### 7.2.6. Área da Planilha

É a área de trabalho do programa, uma grande tabela (estrutura dividida em linhas e colunas) que permite a inserção de dados pelo usuário.

|   | A | B | C | D | E | F | G |
|---|---|---|---|---|---|---|---|
| 1 |   |   |   |   |   |   |   |
| 2 |   |   |   |   |   |   |   |
| 3 |   |   |   |   |   |   |   |
| 4 |   |   |   |   |   |   |   |
| 5 |   |   |   |   |   |   |   |
| 6 |   |   |   |   |   |   |   |
| 7 |   |   |   |   |   |   |   |
| 8 |   |   |   |   |   |   |   |
| 9 |   |   |   |   |   |   |   |

*Figura 7.10 – Parte da área da planilha.*

#### 7.2.6.1. Limites da planilha do Excel 2010

Uma planilha do Excel 2010 possui 1.024.576 linhas (que vão, claro, da linha 1 até a linha 1048576). Repito: um milhão, vinte e quatro mil, quinhentas e setenta e seis linhas em uma planilha do Excel!

As linhas (componentes horizontais da planilha) são, por óbvio, representadas por números!

As colunas, por sua vez, que são as componentes verticais (dividem a planilha lado a lado), são representadas por letras (às vezes uma só, às vezes duas ou três letras).

Uma planilha do Excel oferece 65.536 (sessenta e cinco mil, quinhentas e trinta e seis) colunas. Elas vão da coluna “A” até a coluna “XFD”.

“Ô João, explica essa estória de ‘XFD’ direito, vai... Como é?”

Simples, amigo leitor: as colunas começam sendo representadas com uma letra para cada coluna. Mas isso vai de “A” até “Z”. Depois do “Z”, é necessário começar a contar com duas letras (“AA”, “AB”, “AC”, “AD” etc.).

Isso vai até “ZZ”, pois, ao chegar aí, não tem mais para onde ir com apenas duas letras! Vamos começar, então, a contar com três letras (“AAA”, “AAB”, “AAC” e assim sucessivamente... até “XFD”). Dá uma “olhada”:



|   | X | Y | Z | AA | AB | AC | AD | AE | A |
|---|---|---|---|----|----|----|----|----|---|
| 1 |   |   |   |    |    |    |    |    |   |
| 2 |   |   |   |    |    |    |    |    |   |
| 3 |   |   |   |    |    |    |    |    |   |
| 4 |   |   |   |    |    |    |    |    |   |
| 5 |   |   |   |    |    |    |    |    |   |
| 6 |   |   |   |    |    |    |    |    |   |
| 7 |   |   |   |    |    |    |    |    |   |

*Figura 7.11 – Note a transição da “Z” para a coluna “AA”.*

As linhas (representadas por números) e as colunas (representadas por letras) se encontram em retângulos únicos chamados **Células**. Cada célula é o espaço onde, justamente, podemos colocar nosso conteúdo (textos, números e fórmulas).

Cada célula possui seu endereço próprio, formado pelo endereço da coluna, seguido do endereço da linha que a formam. Por exemplo: o encontro da coluna “B” com a linha “7” forma a célula “B7” (não é 7B! A coluna vem antes!)

#### 7.2.6.2. Planilha versus Pasta de Trabalho

É necessário, caro leitor, chegar a este ponto e esclarecer uma coisa (que já foi dita lá no Windows, na parte sobre os tipos de arquivos mais comuns)...

**Planilha** é o nome dado às “folhas” que existem dentro de um arquivo do Excel. Uma planilha é uma dessas folhas, apenas uma dessas tabelas. Muitos acreditam que se chama planilha o arquivo inteiro do Excel (o arquivo salvo).

**Pasta de Trabalho** é o nome dado ao arquivo do Excel, ao conjunto de várias planilhas (inicialmente três, como vimos). Então, quando se salva um arquivo do Excel, uma pasta de trabalho é salva, e não meramente uma planilha, como muitos acreditam.

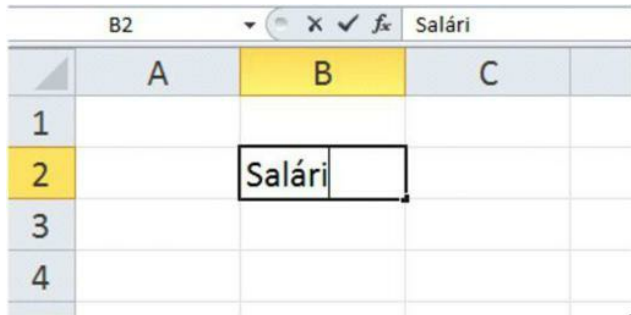
Então, uma pasta de trabalho é um objeto que pode conter várias planilhas. E nunca o contrário!

(Acredite se quiser: isso já caiu em prova! A FCC adora explorar esse conceito!)

#### 7.3. Trabalhando com o Excel

Para inserir dados no Microsoft Excel, basicamente fazemos o seguinte:

1. Selecionar a célula onde iremos escrever (basta clicar na mesma).
2. Digitar o que se deseja.
3. Confirmar a operação (normalmente com ENTER, mas o TAB serve!).



*Figura 7.12 – “Salário” sendo escrito na célula B2.*

Note que o ponto de inserção (a barrinha que fica piscando enquanto você digita) está sendo vista na célula B2. Note também que, à medida que se vai digitando na célula, o conteúdo vai sendo preenchido, também, na barra de fórmulas.

Esse conteúdo ainda não está na célula: ele só ficará lá depois de confirmado (com ENTER ou TAB).

### **7.3.1. Selecionando uma célula**

Para selecionar uma célula, basta clicar nela. Note que uma borda mais escura (chamada borda ativa) indicará que a célula está selecionada. Note também que o nome da célula aparecerá na Caixa de Nome.

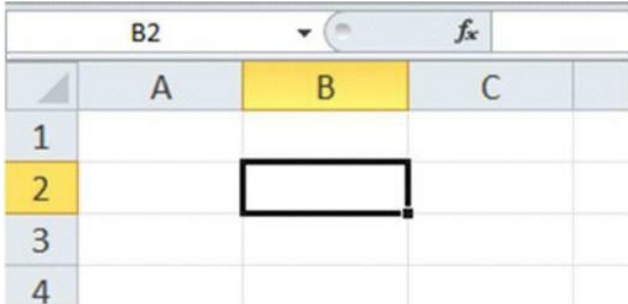


Figura 7.13 – Célula B2 selecionada.

Pode-se selecionar uma célula também usando o teclado. Qualquer uma das teclas a seguir mudará o foco da célula selecionada.

**Setas de direção** (cima, baixo, esquerda e direita): movem a borda ativa, mudando a seleção para as células mais próximas nas respectivas direções para onde apontam.

- **ENTER:** move a borda ativa para a célula abaixo da célula atual. Se o usuário mantiver a tecla SHIFT pressionada, enquanto aciona ENTER (**SHIFT + ENTER**), a borda ativa será movida para a célula acima da atual.

A tecla ENTER é usada para confirmar a escrita de um conteúdo em uma célula, de forma que sempre que esta tecla é pressionada, o conteúdo (se ainda não foi confirmado) será efetivado na célula em que estiver sendo escrito!

- **TAB:** move a borda ativa para a célula à direita da célula atual. Se o usuário acionar SHIFT + TAB, a borda ativa será movida para a célula à esquerda da célula atual.

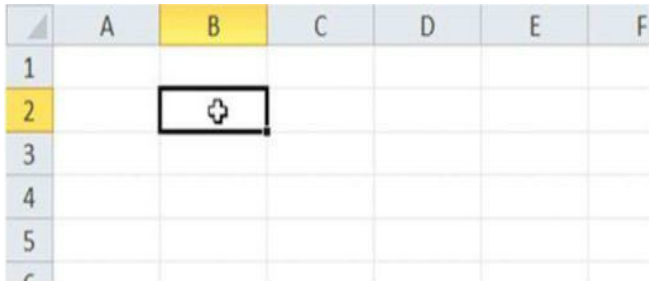
O uso da tecla TAB também serve para confirmar a escrita de um conteúdo em uma célula do Excel.

### 7.3.2. Selecionando várias células

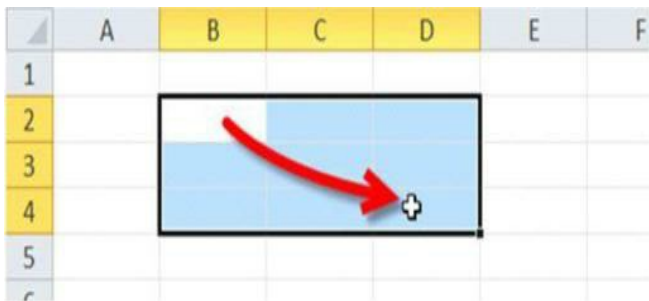
Para selecionar mais de uma célula da planilha, basta usar as teclas **CTRL** e **SHIFT**.

Para selecionar várias células juntas (adjacentes), basta clicar na primeira delas e, segurando SHIFT, clicar na última da sequência. O usuário ainda pode manter a tecla SHIFT pressionada enquanto se move pela planilha (com as setinhas de direção, por exemplo).

O usuário ainda poderá simplesmente arrastar o mouse desde a primeira célula a ser selecionada para a última desejada (figuras seguintes).



*Figura 7.14 – Clica na primeira célula...*



*Figura 7.15 – ... E arrasta o mouse até a última (ou segura SHIFT e clica).*

Para selecionar células não adjacentes (separadas), basta clicar na primeira delas e, pressionando a tecla **CTRL**, clicar nas demais células desejadas.

|   | A | B | C | D | E | F |
|---|---|---|---|---|---|---|
| 1 |   |   |   |   |   |   |
| 2 |   |   |   |   |   |   |
| 3 |   |   |   |   |   |   |
| 4 |   |   |   |   |   |   |
| 5 |   |   |   |   |   |   |
| 6 |   |   |   |   |   |   |

*Figura 7.16 – Seleccionando células não adjacentes.*

Note que quando seleccionamos qualquer conjunto de células, uma delas fica em branco. Não se preocupe, ela também está seleccionada!

Para seleccionar uma coluna inteira da planilha, clique no **cabeçalho da referida coluna** (por exemplo, no retângulo cinza, com a letra B que indica a coluna B). De forma semelhante, para seleccionar uma linha inteira (todas as células de uma linha), pode-se clicar no **cabeçalho da referida linha** (quadrados cinza com os números que ficam à esquerda da planilha).

Para seleccionar todas as células da planilha, basta acionar a combinação de teclas **CTRL + T**. Outra forma de seleccionar toda a planilha é clicar no quadrado cinza que se localiza no topo esquerdo dos cabeçalhos de linha e coluna (entre o cabeçalho da coluna A e o cabeçalho da linha 1). Esse quadradinho é chamado **Cabeçalho da Planilha**.

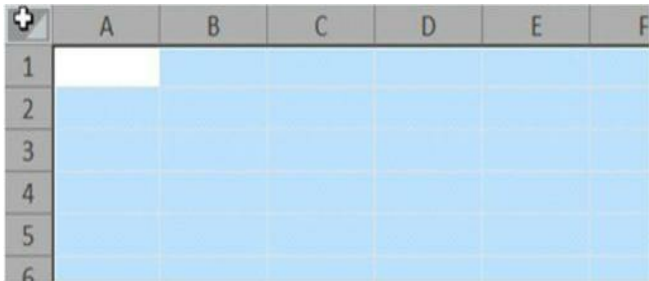


Figura 7.17 – Todas as células selecionadas (note a posição do mouse – a cruz branca).

### 7.3.3. Inserindo dados na planilha

Para inserir qualquer informação na planilha, basta selecionar uma célula qualquer e começar a digitar. Para que o Excel aceite o que foi digitado, o usuário deverá mudar o foco da célula ativa, utilizando uma das formas para mudar a borda ativa de posição (o mais citado é o pressionamento da tecla ENTER).

Caso o usuário, antes de confirmar o conteúdo da célula, pressione a tecla **ESC**, o dado que ele digitou na célula **não será confirmado**, e a célula voltará a apresentar o valor que tinha antes.

Caso o usuário queira **editar** o conteúdo de uma célula previamente preenchida, basta selecionar a referida célula e pressionar a tecla **F2**. A célula irá se “abrir” para o usuário poder modificar seu conteúdo.

Também é possível solicitar a edição da célula aplicando um **duplo clique diretamente nela** ou **um clique na Barra de Fórmulas**.

Caso o usuário deseje apagar o conteúdo inteiro de uma célula, basta selecioná-la e pressionar **DELETE**.

### 7.3.4. Como o Excel entende os dados

Todos os dados que inserimos no Excel são entendidos, pelo programa, de uma dessas três maneiras:

1. Número;
2. Texto;
3. Cálculo.

Se escrevermos um número, o Excel o classificará como tal; se escrevermos algo que não pode ser classificado como número ou cálculo, o Excel o classificará como texto. Seguem alguns exemplos:

- **19:** é um número; **1900** também é número; **1.234,98** idem;
- **Casa:** é um texto;
- **6.5:** também é texto (na configuração do seu computador para informações do Brasil, o número deve ser escrito com vírgula para separar os decimais – portanto, deveria ser **6,5** –; diferente dessa forma de escrever, se torna texto);
- **1,234.98:** é também classificado como texto porque não pode ser classificado como número (simplesmente porque desrespeita as regras sintáticas de escrita de números do Brasil, já que separamos milhares com ponto, e decimais com vírgula – ou seja, exatamente o contrário do que está aí no exemplo).

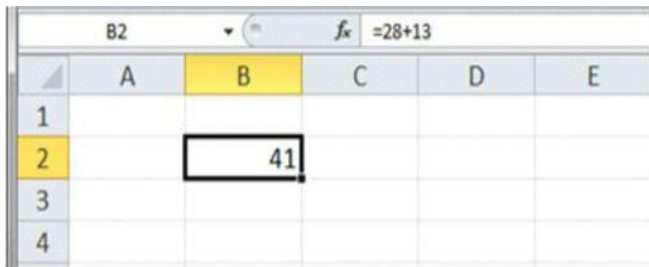
Mas, qual critério o Excel utiliza para classificar o conteúdo das células como cálculos?

#### 7.4. Cálculos – Automatizando o Excel

O Microsoft Excel entende o conteúdo de algumas células como cálculos, e isso faz com que o Excel entenda que precisa executar uma operação antes de mostrar o resultado da célula.

O comportamento padrão do Excel, quando ele encontra um cálculo, é: na hora que o usuário confirma o conteúdo, dando ENTER, por exemplo, o Excel entende o cálculo, executa-o (calcula) e apresenta o resultado dele na célula.

O cálculo, em si, não é mais apresentado na célula, só o resultado. Para ter acesso ao cálculo (expressão) novamente, basta olhar para a barra de fórmulas, ou pedir para editar a célula.



*Figura 7.18 – Cálculo escrito na célula B2 – olha lá em cima!*

Para que o Excel entenda o conteúdo de uma célula como cálculo, basta que o usuário inicie a digitação com um caractere especial, oficialmente, o sinal de “=” (igual).

Mas há mais outros três caracteres que, se inseridos no início da célula, farão o Excel entender o conteúdo como um cálculo: “+” (mais), “-” (menos) e “@” (arroba). O símbolo de @ não é usado para todos os casos, apenas para funções (veremos adiante).

Lembre-se: os cálculos no Excel são entendidos quando se insere, no início da célula, os sinais

de =, +, - e @ (= é o caractere oficial, portanto o mais citado em concursos).

“João, eu poderia escrever ‘+28+13’ e daria no mesmo? E ‘-28+13’ também?”

Veja bem, caro leitor. A expressão =28+13 está somando dois números positivos. A expressão +28+13 também. Logo, elas são idênticas. Já o segundo exemplo que você citou, ou seja, a operação -28+13 está somando um número negativo (-28) com um número positivo (13). O resultado disso será 13-28, ou seja, -15.

“Quer dizer que se eu iniciar o cálculo (fórmula) com o sinal de - (menos), esse sinal tornará negativo o número que o segue imediatamente?”

Sim, sem dúvidas. Começar uma fórmula com o sinal de - (menos) faz com que o sinal em questão se torne não somente um substituto para o = (igual), mas também atua como um “menos unário”, que torna o número à direita dele negativo. Ou seja, a expressão -28+13 é idêntica à expressão =-28+13.

#### 7.4.1. Operando – Operador – Operando

Existe uma regrinha básica para criarmos boas equações (fórmulas) no Excel. Respeitamos a ordem “**Operando – Operador – Operando**”.

“O que é isso, João? Pelo amor de Deus!”

Operando é qualquer número (ou endereço de célula) que sofrerá uma operação aritmética. Operandos são as parcelas da soma ou os fatores de uma multiplicação. São operandos, também, os dividendos e divisores.

Operador é qualquer sinal que realiza operações (os sinais de + e – são exemplos de operadores aritméticos). Existem alguns operadores predefinidos no Excel, como os que veremos a seguir.

Portanto, na expressão =56+78, os números 56 e 78 são operandos. O sinal de “+”, por sua vez, é um operador.

#### 7.4.2. Como fazer cálculos aritméticos

Basta escrever uma equação aritmética com o operador desejado. Verifique a lista dos operadores aritméticos e suas funções:

| Operação      | Operador |  |
|---------------|----------|--|
| Soma          | +        |  |
| Subtração     | -        |  |
| Multiplicação | *        |  |



|             |   |  |
|-------------|---|--|
| Divisão     | / |  |
| Potenciação | ^ |  |
| Porcentagem | % |  |

### 7.4.3. Prioridade dos Operadores

Uma fórmula no Excel pode conter vários operadores aritméticos, como, por exemplo:

$=3*8+10$  (o resultado é 34).

$=3+8*10$  (o resultado é 83 e não 110, como muita gente pensa).

Lembre-se de que o Excel resolverá as operações de uma equação na ordem exigida pela matemática:

1. Potenciação é realizada primeiro;
2. Multiplicação e Divisão são realizadas depois;
3. Adição e Subtração são realizadas por último.

Caso o usuário deseje escrever uma equação que contrarie essa sequência de resolução, poderá alterar a prioridade com o uso de parênteses. Veja exemplos:

$=10+40*10$  resulta em 410

$=(10+40)*10$  resulta em 500

Não há, no Excel, necessidade de usar colchetes ou chaves (na verdade, eles não são aceitos), como fazemos convencionalmente nas equações matemáticas para isolar termos em vários níveis. No Excel usamos somente parênteses. Veja o exemplo:

$=(30*(4+6)+60)/(4*(3+6))$

“E como o Excel resolverá essa equação enorme, João?”

Simples, leitor. O Excel resolverá primeiro o que existe nos grupos de parênteses mais internos e vai seguindo resolvendo os mais externos. Logo ele resolverá assim (olhe a sequência das explicações):

Primeiro, o Excel resolve o que está nos parênteses mais internos, no caso, os trechos  $(4+6)$  e  $(3+6)$ :

$= (30*(4+6)+60)/(4*(3+6)) = (30*10+60)/(4*9)$

Depois, o Excel tenta responder os outros parênteses, mas note que há duas operações no primeiro parênteses:  $(30*10+60)$ . O Excel resolverá a multiplicação primeiro:

$(30*10+60)/(4*9) = (300+60)/(4*9)$

Agora, estará livre para resolver os dois níveis de parênteses:

$(300+60)/(4*9) = 360/36 = 10$

Note bem, caro leitor, este “passo a passo” é mais *SEU* do que do Excel! O Excel tem o jeito dele resolver, e é imediato, automático! A sequência aqui mostrada serve para ajudar você a construir a sua forma de fazer, respeitando as regras, para permitir analisar e propor soluções

para equações grandes que vierem a ser apresentadas em provas!

#### 7.4.4. Referências de Células

Apesar de mostrado nos exemplos anteriores, no Excel raramente usamos somente os valores numéricos dentro dos cálculos aritméticos. É mais comum usarmos referências às células que possuem os valores que apontam para seus endereços na planilha. Veja a seguir:

| B2      fx      =C4+D1 |   |    |    |    |   |  |
|------------------------|---|----|----|----|---|--|
|                        | A | B  | C  | D  | E |  |
| 1                      |   |    |    | 28 |   |  |
| 2                      |   | 40 |    |    |   |  |
| 3                      |   |    |    |    |   |  |
| 4                      |   |    | 12 |    |   |  |
| 5                      |   |    |    |    |   |  |
| 6                      |   |    |    |    |   |  |

Figura 7.19 – Exemplo de cálculo com referências de células (endereços).

Atenção: chamamos de **Referência de Célula** (ou, respeitando a regência, Referência a Células) qualquer dado que, usado em fórmulas, “aponta” para o endereço de uma célula (visando, claro, a seu conteúdo).

Podemos apontar para qualquer célula da planilha, até mesmo para as células que possuem nome amigável (colocado através da Caixa de Nome). Aqui está a razão para usarmos nomes, caro leitor! Podemos usar tais nomes em fórmulas, tornando desnecessário que apontemos para os verdadeiros endereços das células.

|   | B2 |    | $f_x$ | =Valor1+Valor2 |   |  |
|---|----|----|-------|----------------|---|--|
|   | A  | B  | C     | D              | E |  |
| 1 |    |    |       | 28             |   |  |
| 2 |    | 40 |       |                |   |  |
| 3 |    |    |       |                |   |  |
| 4 |    |    | 12    |                |   |  |
| 5 |    |    |       |                |   |  |

Figura 7.20 – A fórmula da soma das células que receberam os nomes Valor1 e Valor2.

Por padrão, o Excel entende uma referência de célula como uma célula da *mesma planilha* em que o cálculo está sendo escrito, a menos que o usuário informe que a célula para a qual a referência aponta está em outra planilha.

“Não entendi esse último parágrafo, João! Explica de novo!”

Claro! O que eu quis dizer foi que se você estiver escrevendo uma fórmula na planilha *Plan1* e essa fórmula é, digamos, **=F9+D8**, o Excel julga que F9 e D8 são células presentes dentro de Plan1. (Porque é a “planilha atual”, ou seja, Plan1 é a planilha onde a fórmula está sendo escrita.) O Excel nunca imaginaria que F9 está em Plan2, porque a fórmula está sendo escrita em Plan1!

Em comparação, é mais ou menos o seguinte: imagina que você mora em São Paulo e que seu telefone é de São Paulo (ou seja, tem código de área 011). Sua irmã também mora em São Paulo e o telefone dela é 3159-0404. Seu irmão mora em Porto Alegre (cujo código de área é 051) e o telefone dele também é 3159-0404.

Pois bem, se você discar do seu telefone o número 3159-0404, sem dizer o código de DDD antes, seu telefone vai ligar para quem? Seu irmão ou sua irmã?

“Claro que ligará para minha irmã, João! Pois eu não disse o DDD 051!”

Isso! Você, quando liga ao telefone, e não especifica o DDD de destino, é entendido como se quisesse ligar para alguém do mesmo código de área! A ideia é a mesma no Excel, entende?

“Mas dá para fazer uma fórmula em Plan1 que aponte para uma célula existente em Plan2, João?”

Sim, leitor! É possível! É muito semelhante a “telefonar para alguém em outra cidade”. É necessário “colocar o código de área antes”.

Para apontar para uma célula em outra planilha, não utilizamos APENAS o endereço da célula. Há uma forma mais específica para apontar: **Planilha!Célula** (Separe o nome da planilha e o nome da célula por um sinal de exclamação).

Ou seja, é como colocar “(0xxYY) Telefone” para discar para o telefone que deseja em outra cidade (com o código YY).

Então, para referir-se a uma célula localizada em uma planilha diferente da planilha onde se está escrevendo a fórmula, deve-se especificar a planilha onde está a célula para a qual se deseja apontar. Para isso devemos utilizar o formato mostrado anteriormente (**Planilha!Célula**).

| B2 |   | fx |    | =D1+Plan3!F9 |   |  |
|----|---|----|----|--------------|---|--|
|    | A | B  | C  | D            | E |  |
| 1  |   |    |    | 28           |   |  |
| 2  |   | 73 |    |              |   |  |
| 3  |   |    |    |              |   |  |
| 4  |   |    | 12 |              |   |  |
| 5  |   |    |    |              |   |  |

Figura 7.21 – Referência para a célula F9 da planilha Plan3 (lá, há o número 45).

“João, mas isso só funciona se a planilha atual e a planilha para onde se vai apontar estejam no **mesmo arquivo**, não é? Dá para fazer uma fórmula que aponte para uma célula presente em outro arquivo salvo no computador?”

Sim! Também é possível apontar para uma célula que esteja em outro arquivo do Excel (outra Pasta de Trabalho). Fazendo outra comparação telefônica, seria como ligar para outro país! Neste caso, você tem que colocar o código do país antes do código da área (algo como +55 011 3159-0404, para ligar para a sua irmã!).

Para isso use a seguinte sintaxe: **[Arquivo]Planilha!Célula** (ao exemplo anterior, apenas adicione, no início, o nome do arquivo entre colchetes). Essa técnica só servirá para o caso de o arquivo denominado **Arquivo** estar salvo dentro da mesma pasta em que o arquivo onde a fórmula está sendo escrita estiver salvo. (Duas pastas de trabalho do Excel salvas na mesma pasta do disco – mesmo diretório.)

|   | SOMA      X ✓ fx      =[Aula 12.xlsx]Plan1!D5 |                        |    |    |   |
|---|-----------------------------------------------|------------------------|----|----|---|
|   | A                                             | B                      | C  | D  | E |
| 1 |                                               |                        |    | 28 |   |
| 2 |                                               | [Aula 12.xlsx]Plan1!D5 |    |    |   |
| 3 |                                               |                        |    |    |   |
| 4 |                                               |                        | 12 |    |   |
| 5 |                                               |                        |    |    |   |

Figura 7.22 – Referência para a célula D5, da planilha Plan1, no arquivo Aula 12.xlsx.

Mas é possível escrever referências de células para arquivos que estão em locais distintos (diretórios, computadores etc.). É possível até apontar para uma célula existente em um arquivo do Excel localizado na Internet.

Note que, para um endereço de arquivo localizado em outro diretório, ou outro computador, como para um site na Internet (um site é apenas outro computador), é preciso lembrar da necessidade do ‘ (apóstrofo), aberto no *início do endereço* e fechado *entre o nome da planilha e o sinal de exclamação*. Outros exemplos de referências feitas a arquivos em outros locais:

=C:\Meus documentos\Projeto\[Orçamento final.xlsx]Plan3!F15

Isso significa: aponte para a célula F15, que está na planilha Plan3, localizada dentro do arquivo Orçamento Final.xlsx, que está salvo na pasta projeto, dentro da pasta Meus Documentos da unidade C: daquele computador. Veja outro:

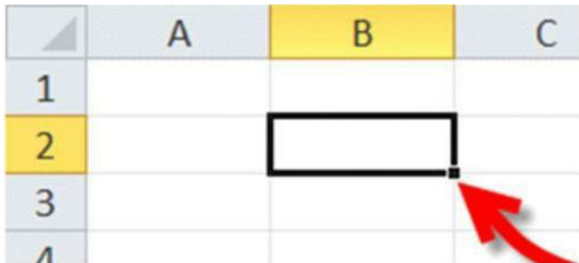
=\\Financeiro1\Planilhas\Pessoal[Ponto.xlsx]Listagem!A11

É uma referência que aponta para a célula A11, dentro da planilha Listagem, que está no arquivo Ponto.xls, gravado na pasta Pessoal, que é subpasta de Planilhas, que está compartilhada a partir do computador chamado Financeiro1 numa rede local. Ufa! (Entendido?)

Além da exata colocação dos apóstrofos (antes do início do endereço e antes do sinal de exclamação), você deve perceber, leitor, que os colchetes envolvem apenas o nome da pasta de trabalho (arquivo do Excel). Só isso!

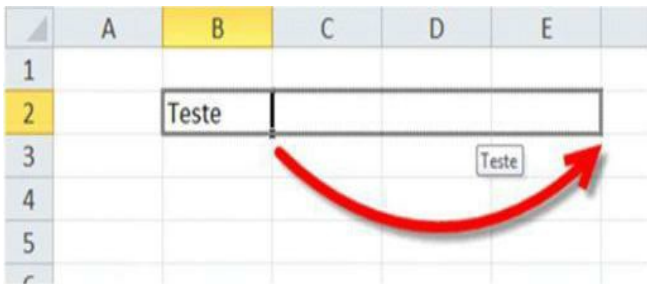
#### 7.4.5. Usando a Alça de Preenchimento

Para facilitar nosso trabalho de preencher a planilha com dados diversos, podemos usar um recurso do Excel chamado alça de preenchimento, que é um pequeno quadrado preto na extremidade inferior direita da célula ativa.



*Figura 7.23 – Alça de preenchimento.*

Como funciona a alça? Basta escrever qualquer valor em uma célula e arrastar pela alça para qualquer direção (acima, abaixo, à direita ou à esquerda). Na maioria dos casos, o Excel irá copiar o valor contido na célula para as demais. Note que a alça não poderá ser arrastada na diagonal.



*Figura 7.24 – Alça sendo arrastada com a palavra “Teste”...*

|   | A | B     | C     | D     | E     |
|---|---|-------|-------|-------|-------|
| 1 |   |       |       |       |       |
| 2 |   | Teste | Teste | Teste | Teste |
| 3 |   |       |       |       |       |
| 4 |   |       |       |       |       |
| 5 |   |       |       |       |       |
| 6 |   |       |       |       |       |

*Figura 7.25 – ... E o resultado disso!*

Em alguns casos específicos, a alça traz resultados muito mais “inteligentes”, como o preenchimento automático de uma sequência de valores predefinidos. As sequências (listas de valores) automáticas no Excel são:

- Meses (por extenso – como “Abril”);
- Meses (abreviados com três letras – como “Jan”);
- Dias da semana (por extenso – como “Domingo”);
- Dias da semana (abreviados com três letras – como “Qui”);

|   | A | B   | C        | D   | E   |
|---|---|-----|----------|-----|-----|
| 1 |   |     |          |     |     |
| 2 |   | Jan | Agosto   | Sex |     |
| 3 |   | Fev | Setembro |     |     |
| 4 |   | Mar | Outubro  |     |     |
| 5 |   | Abr | Novembro |     | Dom |
| 6 |   |     |          |     |     |

*Figura 7.26 – Alça usada em sequências preestabelecidas no Excel.*

#### 7.4.6. Direção e sentido do arrasto em seqüências

Há diferença de comportamento, no Excel, quando se arrasta pela alça de preenchimento, para cima, um valor que faz parte de uma lista conhecida de quando se arrasta esse valor conhecido para baixo.

Quando se arrasta pela alça para *baixo* ou para *a direita*, os valores das seqüências são incrementados a cada célula (ou seja, *Jan* vira *Fev*, que vira *Mar*, que vira *Abr* e assim por diante). Porém, quando a alça é arrastada para *cima* ou para *a esquerda*, os valores são decrementados (diminuídos) a cada célula, o que significa que *Jan* vira *Dez*, que depois vira *Nov*, e assim sucessivamente.

Quando utilizamos a alça com valores de texto que terminam com um número, o Excel também entende que deverá realizar o preenchimento da seqüência. (*Aluno1*, quando arrastado para baixo ou para a direita, virará *Aluno2*, depois *Aluno3*, e assim por diante).

Atenção: se apenas quiser preencher uma seqüência numérica (somente um número na célula), não é suficiente escrever somente um número. (Se assim o fizer, o Excel irá copiar o número em todas as células por onde a alça passou.) Para fazer uma seqüência numérica, o usuário deverá escrever os dois primeiros termos da seqüência (em células adjacentes) e selecioná-los simultaneamente para proceder com o arrasto pela alça.

Ou seja, para obter, como resultado, a seqüência mostrada na figura a seguir, o usuário teve de escrever 1 na célula C1 e 2 na célula C2, depois, selecionou as duas células e procedeu com o arrasto para baixo, criando assim a seqüência mostrada.

|   | A | B | C | D | E |
|---|---|---|---|---|---|
| 1 |   |   | 1 |   |   |
| 2 |   |   | 2 |   |   |
| 3 |   |   | 3 |   |   |
| 4 |   |   | 4 |   |   |
| 5 |   |   | 5 |   |   |
| 6 |   |   |   |   |   |

Figura 7.27 – Alça em seqüências numéricas.



Em alguns casos não é necessário arrastar a alça, apenas aplicar um duplo clique com o mouse na mesma para preencher a sequência desejada. Isso acontece se a coluna imediatamente à esquerda estiver preenchida, e a alça só preencherá até a linha correspondente à última linha preenchida na coluna à esquerda.

|   | A    | B    | C | D | E |
|---|------|------|---|---|---|
| 1 | Blá  | Adm. | 4 |   |   |
| 2 | Blí  | RH   |   |   |   |
| 3 | Bló  | RH   |   |   |   |
| 4 | Blú  | Adm. |   |   |   |
| 5 | Bléu | TI   |   |   |   |
| 6 |      |      |   |   |   |
| 7 |      |      |   |   |   |

Figura 7.28 – Se, desse jeito, o usuário aplicar duplo clique na alça de preenchimento...

|   | A    | B    | C | D | E |
|---|------|------|---|---|---|
| 1 | Blá  | Adm. | 4 |   |   |
| 2 | Blí  | RH   | 4 |   |   |
| 3 | Bló  | RH   | 4 |   |   |
| 4 | Blú  | Adm. | 4 |   |   |
| 5 | Bléu | TI   | 4 |   |   |
| 6 |      |      |   |   |   |
| 7 |      |      |   |   |   |

Figura 7.29 – ... O Excel completará a sequência até encontrar o fim da coluna à esquerda.

#### 7.4.7. A Alça de Preenchimento para fórmulas

Quando utilizamos a alça para preencher células que contenham fórmulas, o Excel realiza uma operação muito interessante. O Excel vai construir, nas demais células, fórmulas com a mesma estrutura da original; porém, com *referências de células atualizadas* de acordo com o movimento realizado a partir da primeira.

|   | A | B | C      | D | E |
|---|---|---|--------|---|---|
| 1 |   |   |        |   |   |
| 2 |   |   |        |   |   |
| 3 |   |   | =F9-D6 |   |   |
| 4 |   |   |        |   |   |
| 5 |   |   |        |   |   |
| 6 |   |   |        |   |   |
| 7 |   |   |        |   |   |

Figura 7.30 – Fórmula “=F9-D6” escrita na célula C3.

|   | A      | B      | C       | D                                                                                   | E      |
|---|--------|--------|---------|-------------------------------------------------------------------------------------|--------|
| 1 |        |        | =F7-D4  |                                                                                     |        |
| 2 |        |        | =F8-D5  |                                                                                     |        |
| 3 | =D9-B6 | =E9-C6 | =F9-D6  | =G9-E6                                                                              | =H9-F6 |
| 4 |        |        | =F10-D7 |  |        |
| 5 |        |        | =F11-D8 |                                                                                     |        |
| 6 |        |        |         |                                                                                     |        |
| 7 |        |        |         |                                                                                     |        |

Figura 7.31 – Resultado do uso da alça de preenchimento.

Se o usuário arrastou a alça para baixo, as fórmulas construídas apresentarão referências para linhas mais baixas (incrementando em um número a cada linha arrastada), ou seja, arrastar B5 para baixo resulta na referência B6. Se caso o usuário arrastar a alça para cima, as referências de células serão atualizadas para uma linha a menos (arrastar B5 para cima resulta na referência B4).

Se o arrasto ocorreu para a esquerda, as próximas fórmulas sofrerão alteração nas referências de colunas, que serão atualizadas para uma letra a menos (ou seja, arrastar B5 para a esquerda cria a referência A5). Por fim, se o arrasto ocorreu para a direita, as referências de colunas das próximas células se apresentarão com uma letra a mais (que significa que arrastar B5 para a direita vai criar a referência C5).

Atenção! As fórmulas não são atualizadas apenas se utilizarmos a alça de preenchimento. Se um usuário escreve uma determinada fórmula usando referências de células e esta for *copiada* (**CTRL + C**), quando colada (**CTRL + V**) em outra célula já *será colada atualizada*.

|   | A | B        | C | D | E |
|---|---|----------|---|---|---|
| 1 |   |          |   |   |   |
| 2 |   | =G12+C10 |   |   |   |
| 3 |   |          |   |   |   |
| 4 |   |          |   |   |   |
| 5 |   |          |   |   |   |
| 6 |   |          |   |   |   |
| 7 |   |          |   |   |   |

Figura 7.32 – A fórmula original foi copiada...

|   | A | B        | C        | D | E |
|---|---|----------|----------|---|---|
| 1 |   |          |          |   |   |
| 2 |   | =G12+C10 |          |   |   |
| 3 |   |          |          |   |   |
| 4 |   |          |          |   |   |
| 5 |   |          | =H15+D13 |   |   |
| 6 |   |          |          |   |   |
| 7 |   |          |          |   |   |

Figura 7.33 – ... e foi colada com as referências atualizadas.

|   | A | B        | C        | D        | E |
|---|---|----------|----------|----------|---|
| 1 |   |          |          |          |   |
| 2 |   | =G12+C10 |          |          |   |
| 3 |   |          |          | =I13+E11 |   |
| 4 |   |          |          |          |   |
| 5 |   |          | =H15+D13 |          |   |
| 6 |   |          |          |          |   |
| 7 |   |          |          |          |   |

Figura 7.34 – A mesma fórmula colada em outra célula.

Atenção: a atualização das referências na fórmula **não ocorre no comando Recortar (CTRL + X)**. Quando usamos esse comando, a fórmula é colada exatamente como estava na célula original, que fica vazia sem a fórmula recortada. (Recortar significa retirar o objeto da origem e colocá-lo apenas onde for colado!)

Lembre-se: só há necessidade de analisar como a fórmula “vai ficar” se o processo descrito

na questão **for copiar** (e posteriormente colar). No caso do ato de **recortar** (**mover** é um sinônimo), a fórmula colada será idêntica à original.

#### 7.4.8. Macete para fórmulas copiadas

Algumas questões de prova pedem que o candidato diga “qual a fórmula que vai aparecer na célula ‘fulana’ se for copiada da célula ‘beltrana’”. Em muitos casos, os candidatos que se veem diante desta questão teimam em querer desenhar a planilha para sair contando quantas linhas e quantas colunas.

“E não é assim, não, João?”

Bem, caro leitor. Assim também funciona, mas leva muito tempo! Que tal se usássemos uma técnica para fazer qualquer tipo de questão dessas com a máxima facilidade? Vamos primeiramente usar uma questão como exemplo:

“(Questão exemplo) Considere que um usuário de computador está editando uma planilha do Microsoft Excel e escreve, na célula B10, a fórmula =G7+D16. Ao copiar essa fórmula, colando-a na célula E6, é correto afirmar que a célula E6 apresentará a fórmula:”

**Primeiro passo:** identifique três informações importantíssimas (sem as quais, a técnica não poderá ser usada, nem a questão respondida).

- **Célula de origem:** B10
- **Célula de destino:** E6
- **Fórmula a ser copiada:** = G7+D16

De posse dessas três informações, siga para o próximo passo.

**Segundo passo:** monte o esquema “**Origem** □ **Destino** :: **Fórmula**”, usando os endereços das células em questão. Ficaria assim:

B10 □ E6 :: =G7+D16

**Terceiro passo:** analise qual foi a mudança da origem para o destino, apenas nos números.

B10 □ E6 :: 10 □ 6 (Mudança: -4)

**Quarto passo:** tendo descoberto a mudança que aconteceu nos números, aplique-a a todos os números da fórmula original.

- **Mudança:** -4
- **Fórmula original (analisando só números):** = \_\_7 + \_\_16
- **Aplicando a mudança:** 7-4 vira 3; 16-4 vira 12)
- **Fórmula atualizada (analisando só números):** = \_\_3 + \_\_12

**Dica:** chegando aqui, descobrimos como ficarão os números na fórmula que será colada em E6. Se, na prova, somente uma das alternativas possuir esses números, é ela a resposta! ;-D

**Quinto passo:** analise qual foi a mudança da origem para o destino, apenas nas letras. A mudança é numérica e está relacionada com a posição das letras no alfabeto.

B10 □ E6 :: B □ E (Mudança: +3, baseado na posição no alfabeto)

**Sexto passo:** tendo descoberto a mudança que aconteceu nas letras da origem para o destino, aplique-a a todas as letras da fórmula original.

- **Mudança:** +3

- **Fórmula original (analisando só letras):** = G\_\_ + D\_\_
- **Aplicando a mudança:** (G+3 vira J; D+3 vira G)
- **Fórmula atualizada (só letras):** = J\_\_ + G\_\_

Atenção: eu sei que “G+3” é meio estranho de ler... Mas entenda: G + três letras (ou seja, “salte” no alfabeto mais três letras após a letra G... o que temos?). Isso dá, depois da “G”, as letras “H”, “I” e, finalmente, “J”.

**Sétimo (e último) passo:** escreva a fórmula resultante com as colunas (letras) e linhas (números) atualizadas. Ou seja, a fórmula resultante é =**J3+G12**.

Em suma, se a fórmula =G7+D16, presente na célula B10, for copiada para a célula E6, ela será reescrita, em E6, como =J3+G12.

“João, esse ‘macete’ me pareceu muito demorado para fazer durante a prova! Lembre-se de que não temos muito tempo!”

Você entendeu errado, leitor! Esse macete é ótimo! É rápido... Especialmente se for escrito (manuscrito) na hora da prova. Veja a questão a seguir:

“(Questão para testar mais ainda) Um usuário do Excel escreveu, na célula G7, a fórmula =F9-H5 e a copiou, colando-a na célula D11. A fórmula que será gravada em D11 é”

O candidato começa escrevendo o esquema no caderno de provas:

$$G7 \rightarrow D11 \quad = F9 - H5$$

*Figura 7.35 – Início da resolução: desenhando o esquema.*

Depois disso, é só analisar quanto variou em números (linhas) e quanto variou em letras (colunas). Dá para fazer os dois ao mesmo tempo, desenhando as linhas mostradas a seguir. (Observe que o candidato está escrevendo no caderno de provas – arrume algum espaço.)



$$= F9-H5$$

Figura 7.36 – Descobrimo as mudanas.

Note que identificar a mudana dos  meros foi f cil (de 7 para 11 d  +4). Mas “bateu uma pequena d vida” em saber quanto havia “mudado” de G at  D. O nosso candidato do exemplo escreveu um “pequeno alfabeto” e saiu “saltando” de letra em letra, do G pro D, para descobrir que houve tr s saltos (ficou -3 porque os saltos foram “voltando” no alfabeto).

Agora o nosso candidato imagin rio vai aplicar as mudanas que descobriu (+4 para os  meros e -3 para as letras). Pode ser que seja necess rio usar o “alfabeto de apoio” novamente. N o tenha vergonha de us -lo!).



$$= \cancel{F9} - \cancel{H5}$$

$$= C13 - E9$$



Figura 7.37 – Resolvendo a quest o.

Pronto! Nosso candidato descobriu que a f rmula =F9-H5, escrita originalmente em G7, ser 

reescrita como =C13-E9 se for copiada e colada em D11. Parece ser um procedimento muito mais rápido, não?

Mas não se esqueça disto: se o enunciado falar em “recortar” ou “mover”, não há necessidade de fazer esses cálculos que fizemos, pois a fórmula simplesmente não será alterada.

Vamos partir agora para um novo assunto: as referências absolutas.

#### 7.4.9. Usando referências absolutas

Chamamos de referência absoluta (ou fixa) a referência de célula que não se altera com o uso da alça de preenchimento ou com os comandos copiar/colar.

Em certos casos, é necessário que uma referência de célula não se altere durante o arrasto com a alça ou durante os comandos copiar/colar. (Isso depende, da estrutura da planilha em questão.) Para fixar uma referência, basta colocar um **\$ (cifrão)** imediatamente antes da parte da referência que se deseja fixar.

Exemplos:

=C9 (C livre; 9 livre)

=C\$9 (C livre; 9 fixo)

=\$C9 (C fixo; 9 livre)

=\$C\$9 (C fixo; 9 fixo)

Dizemos que a referência que não possui cifrão é relativa (a primeira da listagem anterior); uma referência que possui as duas partes com cifrão é chamada referência absoluta (a última do exemplo anterior); e quando uma referência possui apenas um componente fixo (linha, como no segundo exemplo ou coluna, no terceiro), é chamada referência mista.

Alguns autores, porém (e alguns elaboradores de provas, também), gostam de dizer que C9 é composto por duas referências: “C” é a referência de coluna, e “9” é a referência da linha. Desta forma, que possuir “\$” é fixo, e quem não o possuir, é relativo.

Ou seja, pode ser que a expressão **C\$9** seja lida, na sua prova, como sendo “referência relativa de coluna e absoluta de linha”. Já viu, né?

Veja, nas figuras a seguir, exemplos práticos do funcionamento deste recurso.



|   | A | B        | C                                                                                 | D |
|---|---|----------|-----------------------------------------------------------------------------------|---|
| 1 |   |          |                                                                                   |   |
| 2 |   | =C\$9*F2 |                                                                                   |   |
| 3 |   | =C\$9*F3 |                                                                                   |   |
| 4 |   | =C\$9*F4 |                                                                                   |   |
| 5 |   | =C\$9*F5 |                                                                                   |   |
| 6 |   | =C\$9*F6 |                                                                                   |   |
| 7 |   |          |  |   |

*Figura 7.38 – A fórmula original “=C\$9\*F2” e o resultado do arrasto com a alça de preenchimento.*

No exemplo acima, note que a referência à linha 9 não variou durante o arrasto!

|   | A | B          | C                                                                                   | D         |  |
|---|---|------------|-------------------------------------------------------------------------------------|-----------|--|
| 1 |   |            |                                                                                     |           |  |
| 2 |   | = \$C9*F2  | = \$C9*G2                                                                           | = \$C9*H2 |  |
| 3 |   | = \$C10*F3 |                                                                                     |           |  |
| 4 |   | = \$C11*F4 |                                                                                     |           |  |
| 5 |   | = \$C12*F5 |                                                                                     |           |  |
| 6 |   | = \$C13*F6 |                                                                                     |           |  |
| 7 |   |            |  |           |  |

*Figura 7.39 – A fórmula “= \$C9\*F2” e as resultantes do arrasto da alça.*

Note que a referência à coluna C não variou durante o arrasto.

A referência antecedida do \$ não vai variar mesmo se o usuário usar os comandos copiar/colar, como vemos a seguir.

|   | A | B          | C          | D          | E          |
|---|---|------------|------------|------------|------------|
| 1 |   |            |            | =SC10+G\$4 |            |
| 2 |   | =SC11+E\$4 |            |            |            |
| 3 |   |            |            | =SC12+G\$4 |            |
| 4 |   |            |            |            |            |
| 5 |   |            | =SC14+F\$4 |            | =SC14+H\$4 |
| 6 |   |            |            |            |            |
| 7 |   |            |            |            |            |

Figura 7.40 – A fórmula original “=SC11+E\$4” foi copiada e colada várias vezes.

(O C da primeira referência e o 4 da segunda não variam.)

#### 7.4.9.1. Usando F4 para construir referências absolutas

Caso o usuário queira uma forma fácil de colocar os \$ nas referências, aqui vai uma dica: escreva, por exemplo, a fórmula =D8 e, com o cursor ainda encostado no “8” da referência, pressione a tecla F4. Você verá que o D8 virará \$D\$8 e, se pressionar F4 novamente, será alternado entre \$D8, D\$8, D8 e \$D\$8... É muito legal!

#### 7.4.10. Macete na hora da prova para referências absolutas

Com esta dica, você ganhará um tempo importante! Não tenha dúvidas! Vamos analisar a questão a seguir:

“(Questão para testar \$) Um usuário do Excel copiou a fórmula =C\$8-\$D7 que estava na célula B11 para a célula G15. A fórmula que será apresentada nesta última célula será:”

- a) =H\$12-\$I11;
- b) =H\$12-\$D7;
- c) =H\$8-\$D11;
- d) =H\$7-\$I11;
- e) =C\$12-\$H7.

E então, caro leitor? Já respondeu? Já sabe qual a resposta?

“Peraí, João... Deixa eu fazer o cálculo...”

Que nada, deixa não! Do que você precisa? Precisa comparar, apenas! Compare os valores

fixos. Compare as referências que estão presas pelo \$. A resposta não pode ter valores diferentes daqueles presos por \$ no enunciado! (Lembre-se de que o \$ fixará aquela referência!)

A fórmula do enunciado é =C\$8-\$D7, logo, qualquer fórmula derivada desta (por arrasto da alça ou por cópia/colagem) será =\_\_\$8-\$D\_\_ (as partes fixas nunca vão se alterar). Portanto, a única resposta possível à questão é a “(C) =H\$8-\$D11”.

(Tudo bem que eu escrevi uma questão que apresenta apenas uma alternativa com os valores das referências mantidos, mas as bancas fazem isso às vezes!)

Se não sobrar apenas uma única resposta, pode ser que sobre duas. Isso significa que você terá “cortado da sua vida” três alternativas das cinco. É uma ajuda e tanto!

Atenção: se o elaborador da questão quiser realmente dar trabalho, ele vai criar as cinco alternativas com os mesmos valores no \$. Aí cabe a você, caro leitor, achar a resposta por meio daquela técnica que descrevi anteriormente. (Observe que a técnica só será usada nas referências que não estão fixas pelo \$.)

#### 7.4.11. Usando as funções do Excel

Funções são comandos que acompanham o programa Excel para facilitar nosso trabalho em relação a alguns cálculos específicos. As funções, na verdade, realizam cálculos predefinidos. O Excel possui cerca de 400 funções, para as mais variadas finalidades, desde matemática e trigonometria até matemática financeira e estatística.

Atenção: toda função apresenta um resultado, ou, como costumamos chamar, retorna um resultado (pode ser que o pessoal da prova use esse termo). Ou seja, para o “informatiquês”, **retornar** significa **resultar**. Ou seja, se uma função “retorna 30” é porque simplesmente o resultado daquela função é 30.

Toda função do Excel pode ser solicitada da seguinte forma:

=**NOME**(**ARGUMENTOS**)

onde:

**NOME** é o nome da função (o usuário deve saber o nome da função que deseja utilizar, isso é mais que óbvio);

**ARGUMENTOS** são informações que precisam ser dadas à função para que ela proceda com o cálculo e nos traga o resultado desejado.

Se o usuário precisar informar mais de um argumento à função, pode separá-los, dentro dos parênteses, pelo sinal de ponto e vírgula.

Assim:

=**NOME**(**ARGUMENTO1**; **ARGUMENTO2**)

Algumas funções do Excel são o que eu chamo de “funções intransitivas” (fazendo alusão aos verbos intransitivos, que não requerem objetos). São funções que, quando escritas, não precisam de argumentos, ou seja, não pedem nenhum argumento como complemento.

Uma função que não pede argumentos é escrita assim: =**NOME**() (é necessário abrir e fechar os parênteses). Eis algumas funções “intransitivas”, ou seja, que não pedem argumentos para trazer respostas.

Note bem, caro leitor! Eu chamo as funções de “intransitivas” e “transitivas” para fazer comparação com os verbos (estudo de sintaxe da língua portuguesa). Ou seja, estou apenas

“comparando” com verbos que precisam de complemento (transitivos) e os que não precisam (intransitivos).

Oficialmente, no Excel, esses nomes (transitivas e intransitivas) não são usados!

#### 7.4.11.1. Funções “Intransitivas”

##### =HOJE()

Esta função retorna (traz como resultado) a data atual do seu computador. Se o calendário do seu computador, caro leitor, estiver mal configurado, esta função retornará a data que ele estiver marcando, mesmo que não seja a data real.

O resultado é apresentado no formato DD/MM/AAAA, como em 12/05/2013.

##### =AGORA()

Retorna a data e a hora atuais, no formato DD/MM/AAAA HH:MM, como em *12/05/2013 19:19*.

##### =PI()

Retorna 3,141592654 (o valor da constante trigonométrica  $\pi$  com nove casas decimais).

##### =ALEATÓRIO()

Retorna um número aleatório maior ou igual a 0 (zero) e menor que 1. É uma função muito comum usada em sorteios (eu mesmo a utilizo muito quando faço sorteios nas salas de aula).

#### 7.4.11.2. Funções “Transitivas”

Além dessas funções, há outras que chamo de “funções transitivas” (pedem apenas um argumento). Essas funções podem ser vistas a seguir.

##### =RAIZ(Núm)

Retorna o valor da raiz quadrada do número descrito como argumento. Ou seja, se o usuário escrever =RAIZ(144), o resultado apresentado será 12.

O argumento da função RAIZ só pode ser um número (ou, é claro, uma referência que aponte para uma célula contendo um número).

##### =TRUNCAR(Núm)

A função TRUNCAR é usada para cortar as casas decimais de um número, retornando, neste caso, apenas a parte inteira deste. Quando usamos apenas um argumento (o número a ser truncado), o resultado é o valor do número inteiro (ou seja, a desconsideração total da parte decimal).

Ou seja, ao utilizar =TRUNCAR(3,9832), o resultado obtido é 3.

A função TRUNCAR também pode ser usada de forma politransitiva (se oferecermos mais de um argumento), mas veremos isso mais adiante.

##### =ABS(Núm)

Retorna o valor absoluto (sem sinal) de um número. Ou seja, =ABS(-30) resulta em 30.

### **=COS(Núm)**

Retorna o cosseno de um ângulo (esse “Núm” no argumento representa um ângulo em radianos).

### **=COSH(Núm)**

Retorna o cosseno hiperbólico de um número (“núm” é um número real qualquer).

### **=SEN(Núm)**

Retorna o seno de um ângulo (“núm” é um ângulo em radianos).

### **=SENH(Núm)**

Retorna o seno hiperbólico de um número (“núm” é um número real qualquer).

### **=TAN(Núm)**

Retorna a tangente de um ângulo (“núm” é um ângulo em radianos).

### **=TANH(Núm)**

Retorna a tangente hiperbólica de um número (“núm” é um número real qualquer).

## **7.4.11.3. Funções “Politransitivas”**

Há, também, algumas funções que exigem mais argumentos: são as funções que apelidei carinhosamente de “politransitivas”. Note novamente, caro leitor, que esses termos que se referem à transitividade são meramente “criações minhas”, por isso não têm como aparecer em prova.

Sempre que uma função é escrita contendo mais de um argumento, é necessário separá-los (os argumentos) por meio do sinal de “,” (ponto e vírgula).

Vamos a algumas das mais comuns funções “politransitivas” do Excel. Essas, sim, a despeito das anteriores que mostrei, são muito comuns em provas de diversas bancas examinadoras:

### **=SOMA(Núm1; Núm2; Núm3...)**

A função SOMA simplesmente realiza, como você já deve ter deduzido, o somatório dos números descritos no argumento. Uma função =SOMA(34;11;45) resulta em 90.

### **=MÉDIA(Núm1; Núm2; Núm3...)**

Retorna a média aritmética dos valores descritos no argumento. Uma função =MÉDIA(34;11;45) resulta em 30.

Note que a função MÉDIA vai somar os vários argumentos e em seguida dividir o resultado dessa soma pela quantidade de argumentos existentes.

### **=MÁXIMO(Núm1; Núm2; Núm3...)**

A função MÁXIMO simplesmente retorna o maior número encontrado dentre os argumentos. Uma função =MÁXIMO(34;11;45) resultaria em 45.

### **=MÍNIMO(Núm1; Núm2; Núm3...)**

A função MÍNIMO é o que chamo de “irmã mais nova” da função MÁXIMO. Ela retorna o

menor número encontrado dentre os argumentos. (Creio que isso você já havia concluído sozinho, não é?) Uma função =MÍNIMO(34;11;45) resultaria em 11.

**=MULT(Núm1; Núm2; Núm3...)**

A função MULT retorna o produto (multiplicação) dos números descritos no argumento. Uma função =MULT(34;11;45) resulta em 16830.

### **Mais Sobre as Funções**

“Ei, João, as funções só podem ser escritas com, no máximo, três argumentos? Você usou apenas três argumentos em todas elas!”

Não, leitor! Posso usar mais argumentos, como em =SOMA(34;23;12;90;120;34567;2;45). Na verdade, na função SOMA e nas demais vistas anteriormente, podemos usar até 255 argumentos separados por ponto e vírgula.

“Só posso usar números nos argumentos?”

Não! Podemos usar também referências de células, desde que elas apontem para células que contenham valores numéricos. Afinal, leitor, há de convir que é meio estranho pedir ao Excel que calcule a média entre “23” e “casa”, não é? Portanto, uma função dessas que vimos pode ser perfeitamente escrita assim:

=SOMA(B4;B5;B10;C8;C14;23;D9)

Note que há um número no meio dos argumentos (23, em meio a tantas referências). Não há nenhum problema nisso! O Excel aceitará normalmente essa função. Ele buscará os valores existentes nas células B4, B5, B10, C8, C14 e D9, e somará todos eles ao valor já conhecido 23.

Portanto, nas funções que acabamos de ver, podemos escrever, como argumentos, tanto números (valores numéricos) como endereços (referências) de células! Mas... tem mais!

Essas funções simplesmente ignoram células que contenham textos, se assim forem apontadas para elas. Ou seja, se qualquer uma das funções descritas neste tópico (MÁXIMO, MÍNIMO, MÉDIA, SOMA e MULT) for apontada para uma célula vazia ou para uma célula que contenha texto, ela simplesmente vai ignorar a célula para fins de cálculo.

Isso significa que nem a função média irá considerar a célula em questão na hora de “dividir” pelo número de itens (ou seja, se a MÉDIA for apontada para cinco células no intervalo e uma delas estiver vazia, ou com texto, a MÉDIA irá dividir por 4).

#### **7.4.12. Usando intervalos de células**

Em alguns casos, usar várias células pode ser bastante difícil (por exemplo, quando se tem de apontar para dezenas de endereços, como no exemplo a seguir.). Já imaginou se o usuário precisar somar todas as células existentes de E1 até E10? Uma forma seria fazer:

=SOMA(E1;E2;E3;E4;E5;E6;E7;E8;E9;E10)

Preste atenção, porém, no fato de que essa função está apontando para diversas células sequencialmente dispostas (uma vizinha à outra). Isso cria um *intervalo de células*.

Intervalo (ou Matriz) de células é um conjunto ininterrupto de células dispostas de forma adjacente (células vizinhas). Para apontar para as várias células da função mostrada, usa-se a seguinte sintaxe:

=SOMA(E1:E10)

Um intervalo de células é escrito com o uso do  **sinal de dois pontos**  entre a referência inicial e a final do intervalo. Para citar um intervalo de células, não é necessário citar mais nenhuma outra célula além da primeira e da última.

Para ficar fácil de entender, lembre-se disto: o sinal de “;” (ponto e vírgula) é usado para separar argumentos em uma função e pode ser lido como “e”. O sinal de “:” (dois-pontos) é usado para informar ao Excel sobre um intervalo de células e pode ser lido como “até”.

Então, =SOMA(B2:B10) pode ser lido como “Realize a SOMA de B2 até B10”.

E a função =SOMA(B2;B3;B5;B10) pode ser lida como “Realize a SOMA entre B2 e B3 e B5 e B10”.

É possível fazer intervalos verticais, horizontais e diagonais. Os intervalos verticais são aqueles em que se citam células adjacentes pertencentes a apenas uma coluna. Os intervalos horizontais, por sua vez, são aqueles que citam apenas células de uma mesma linha.

Os intervalos diagonais citam células em linhas e colunas diferentes, gerando, como resultado, um conjunto de células que forma um retângulo.

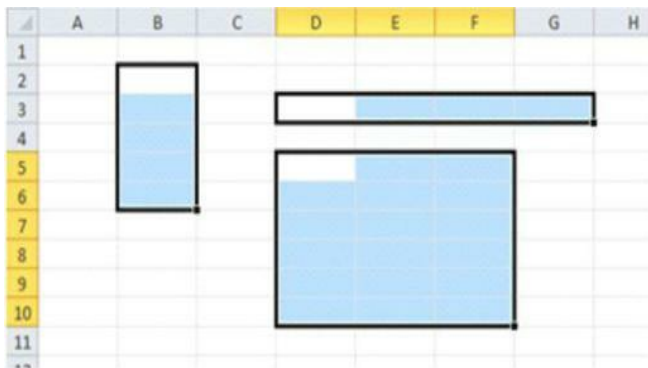


Figura 7.41 – Exemplo de três intervalos.

Observe na figura anterior os três intervalos:

- **Vertical (B2:B6):** envolve as células B2, B3, B4, B5 e B6;
- **Horizontal (D3:G3):** envolve as células D3, E3, F3 e G3;
- **Diagonal (D5:F10):** envolve as células D5, D6, D7, D8, D9, D10, E5, E6, E7, E8, E9, E10, F5, F6, F7, F8, F9 e F10.

Observe que nos intervalos só é necessário citar a primeira e a última célula e que no caso do

intervalo diagonal, a primeira célula é a mais superior e mais à esquerda (no nosso caso, D5), e a última célula do intervalo é a mais inferior e mais à direita (no caso, F10).

Ah, quase ia me esquecendo: também podemos intercalar dois-pontos com ponto e vírgulas para obter interessantes argumentos para nossas funções, como a seguir:

=SOMA(B2:B30;D2:D30)

Essa função irá somar os valores da coluna formada por B2 até B30 com os valores da coluna de D2 até D30. Observando a figura a seguir, podemos ter uma ideia mais clara a respeito das funções e de várias maneiras de usá-las.

|   | A                                                     | B           | C           | D           |  |
|---|-------------------------------------------------------|-------------|-------------|-------------|--|
| 1 | <b>Candidatos a Cargos Públicos (Fiscais) no País</b> |             |             |             |  |
| 2 |                                                       | <b>2010</b> | <b>2011</b> | <b>2012</b> |  |
| 3 | <b>Nordeste</b>                                       | 50000       | 60000       | 80000       |  |
| 4 | <b>Norte</b>                                          | 20000       | 10000       | 15000       |  |
| 5 | <b>Sul</b>                                            | 35000       | 70000       | 56000       |  |
| 6 | <b>Sudeste</b>                                        | 122000      | 110000      | 90000       |  |
| 7 | <b>Centro-Oeste</b>                                   | 30000       | 40000       | 50000       |  |
| 8 |                                                       |             |             |             |  |
| 9 |                                                       |             |             |             |  |

Figura 7.42 – Exemplo de planilha.

A função =SOMA(B3:C7) totaliza o número dos candidatos a cargos fiscais nos anos de 2010 e 2011 em todas as regiões do país.

A função =MÉDIA(B7:D7) calcula a média anual de candidatos a cargos fiscais na região Centro-Oeste (considerando os três anos mostrados).

Para saber qual foi o total de candidatos no Nordeste entre 2010 e 2012, usa-se =SOMA(B3:D3).

#### 7.4.12.1. Alguns “segredos” dos intervalos de células

Aqui vão alguns “alertas” interessantes sobre os intervalos de células:



## Vários Pontos

Não é somente o símbolo de “.” que pode ser usado para indicar intervalos. Você também poderá usar um “.” (ponto), ou “..” (dois pontos seguidos), ou “...” (três pontos seguidos), ou “.....” (“n” pontos seguidos).

Ou seja, escrever

=SOMA(B1.B5) ou =SOMA(B1..B5) ou =SOMA(B1...B5) ou =SOMA(B1.....B5)

É o mesmo que escrever =SOMA(B1:B5).

Ahhh, claro, o Excel irá transformar qualquer função que usou “....” em “.” automaticamente quando você acionar o ENTER!

## Espaço

Você poderá ser questionado, em qualquer prova, acerca de alguma função do seguinte tipo:

=SOMA(C1:D9 A3:F6)

Note que há um espaço (conseguido, claro, pela barra de espaço do seu teclado) entre os dois intervalos. O espaço entre intervalos no Excel tem o objetivo de atuar como ***Operador de Interseção***.

Isso significa que esse sinal (o espaço) é usado para que se considere, apenas, a interseção entre os dois intervalos (ou seja, só se vão considerar as células que fazem parte dos dois intervalos).

|    | A  | B  | C  | D  | E                  | F  |
|----|----|----|----|----|--------------------|----|
| 1  | 12 | 15 | 60 | 29 | 29                 | 15 |
| 2  | 20 | 47 | 30 | 47 | 32                 | 49 |
| 3  | 34 | 59 | 48 | 40 | 22                 | 17 |
| 4  | 18 | 49 | 21 | 34 | 13                 | 26 |
| 5  | 10 | 20 | 46 | 13 | 17                 | 40 |
| 6  | 50 | 21 | 35 | 24 | 22                 | 12 |
| 7  | 30 | 56 | 24 | 16 | 12                 | 31 |
| 8  | 29 | 31 | 59 | 60 | 10                 | 27 |
| 9  | 18 | 55 | 13 | 12 | 51                 | 55 |
| 10 |    |    |    |    |                    |    |
| 11 |    |    |    |    | =SOMA(C1:D9 A3:F6) |    |

Figura 7.43 – Interseção entre dois intervalos.

No exemplo, a função Soma resultará em 261, pois somará apenas o que estiver nas células C3 até D6 (borda mais grossa mostrada na figura anterior), pois estas células são as células que pertencem aos dois intervalos mostrados.

### Funções que Usam Intervalos

Todas as funções que vimos como “politransitivas”, até o presente momento, admitem tanto o uso de “;” (ponto e vírgula) – para indicar uma separação entre os argumentos da função – como admitem “:” (dois-pontos) para indicar intervalos.

Claro que estas funções (SOMA, MÁXIMO, MÍNIMO, MÉDIA, MULT) também admitem o uso do espaço, para fins de apontar interseções de intervalos.

Portanto, nestas funções, há uma infinidade de combinações possíveis de referências de células, que podem, ou não, ser cobradas nas próximas provas que você enfrentar!

Exemplos? Claro! Aqui vão:

=SOMA(30;C2:C8;D1:F8 A4:H6)

=MÉDIA(B10:B20;C1:C80;D1:D30)

=MÁXIMO(C1:C90 D3:F20;I12)

E assim por diante...

Não se esqueça disso! Todas as funções que nós vimos (as politransitivas) admitem receber, como argumentos: números, endereços de células, intervalos e interseções de intervalos!

#### 7.4.13. Expressões mais complexas

Algumas bancas examinadoras gostam de “mexer mais embaixo” no que diz respeito às funções do Excel. São cobradas, muitas vezes, funções muito grandes com vários argumentos diversos que, inclusive, se misturam com cálculos aritméticos variados. Uma verdadeira “festa”!

Vamos tomar a planilha a seguir como base para os nossos exemplos futuros.

|    | A  | B  | C  | D  | E   | F  | G |
|----|----|----|----|----|-----|----|---|
| 1  | 15 | 85 | 35 | 85 | 65  | 85 |   |
| 2  | 70 | 40 | 5  | 60 | 60  | 80 |   |
| 3  | 50 | 95 | 25 | 55 | 100 | 10 |   |
| 4  | 50 | 20 | 90 | 35 | 75  | 10 |   |
| 5  | 35 | 90 | 30 | 40 | 40  | 55 |   |
| 6  | 90 | 70 | 15 | 65 | 20  | 50 |   |
| 7  | 55 | 25 | 90 | 95 | 80  | 80 |   |
| 8  | 35 | 30 | 55 | 65 | 90  | 35 |   |
| 9  | 40 | 50 | 15 | 75 | 30  | 25 |   |
| 10 | 35 | 70 | 50 | 30 | 75  | 95 |   |
| 11 |    |    |    |    |     |    |   |

Figura 7.44 – Planilha de exemplo.

Agora vamos analisar uma função que pode muito bem ser apresentada para você, caro leitor: =MÉDIA(SOMA(B2;B5);MULT(A1;C6)+F3;(C9/C2)\*D1; MÁXIMO (E3:E5)\*3)

“Eita, João... Vai com calma... Pegou pesado agora... É assim que cai, é?”

Caro leitor, é quase isso mesmo! Em algumas provas, especialmente as da Cesgranrio e Fundação Getúlio Vargas (FGV), as questões de Excel exigem um raciocínio muito bom e uma atenção extrema.

Felizmente, vou mostrar um jeito que poderá facilitar a sua vida nesses casos:

1. Em primeiro lugar, tente entender quem é a “função principal” e quantos argumentos ela tem.

De preferência, substitua os argumentos corretos por letras ou indicadores mais fáceis de

escrever. Vejamos: a função “mais abrangente” é a MÉDIA, afinal, todas as demais estão dentro dela.

Note que há quatro argumentos (pois há três sinais de “,” relativos à função MÉDIA) e isso me permite reescrever a função desta forma:

=MÉDIA(a; b; c; d)

Sendo que,

a = SOMA(B2;B5);

b = MULT(A1;C6)+F3;

c = (C9/C2)\*D1; e

d = MÁXIMO(E3:E5)\*3

2. Depois de separar os argumentos da função principal em “equações” individuais, agora chegou a hora de resolvê-las (também individualmente):

a = SOMA(B2;B5) = B2 + B5 = 40 + 90 = **110**

b = MULT(A1;C6)+F3 = A1 \* C6 + F3 = 15 \* 15 + 10 = **235**

c = (C9/C2)\*D1 = (15 / 5) \* 85 = 3 \* 85 = **255**

d = MÁXIMO(E3:E5)\*3 = MÁXIMO(100; 75; 40) \* 3 = 100 \* 3 = **300**

3. Pronto! Agora é só usar os resultados que achamos para as equações a, b, c e d e aplicá-los na função principal.

=MÉDIA(a; b; c; d) será escrita como =MÉDIA(110;235;255;300). Isso vai resultar, na média aritmética entre esses 4 números, que é **230**.

Viu como é tranquilo, amigo leitor?!

#### 7.4.14. Usando funções menos comuns

Dentro da grande quantidade de funções que o programa apresenta, existe a possibilidade de nos depararmos, em concursos, com algumas funções incomuns no dia a dia. Portanto, seguem algumas das funções que não são tão facilmente usadas em nosso cotidiano.

##### 7.4.14.1. Funções de Contagem

=CONT.VALORES(Célula1; Célula2; Célula3...)

Essa função retorna quantas células, dentro das que forem indicadas, não estão vazias. Esta função pode ser utilizada, também, apontando para intervalos.

E, claro, se pode ser apontada para um intervalo, pode ser apontada para uma interseção de intervalos, também!

Veja um exemplo:

=CONT.VALORES(B1:B10)

Não se esqueça disso! A função CONT.VALORES serve para contar células que possuam algum conteúdo (qualquer conteúdo é válido: números, textos etc.). Essa função somente NÃO CONTA células vazias!

=CONT.NÚM(Célula1;Célula2;Célula3...)

Essa função conta quantas células, dentre as indicadas, são formadas por números (ou seja, na

contagem, essa função ignora as células que contêm texto e as células vazias).

Esta função, óbvio, também aceita ser escrita por intervalos e interseções de intervalos.

Já deu para perceber a essa altura, não é, caro leitor? Se uma função pode ser usada indicando várias células, então ela pode ser usada por meio da indicação de um intervalo!

E, claro, se podemos indicar intervalos, podemos indicar a interseção de intervalos!

Exemplo: =**CONT.NÚM(B2:B15)**

### =**CONT.SE(Intervado;Critério)**

Essa função conta quantas vezes um determinado valor (número ou texto) aparece em um intervalo de células (o usuário tem de indicar qual é o critério a ser contado).

Diferentemente das anteriores, a função CONT.SE não pode ser escrita por meio de várias células (separadas por ponto e vírgula). Aqui, necessariamente, deve-se indicar um intervalo.

O ponto e vírgula vai ser usado para separar o intervalo de busca do critério a ser pesquisado.

Exemplo: =**CONT.SE(B2:B15; “Teste”)**

Nesse exemplo, o Excel irá contar quantas células possuem o valor Teste dentro do intervalo de B2 até B15.

Se você quiser, caro leitor, contar, por exemplo, quantas vezes o número 38 aparece no intervalo que vai da célula B10 até a célula B200, é só escrever assim:

=**CONT.SE(B10:B200;38)**

Se o critério a ser pesquisado é um número (38, no caso), não precisa colocá-lo entre aspas (mas, se o fizer, dá no mesmo). Só é necessário colocar entre aspas se o argumento em questão for um texto (ou uma “expressão lógica”, como veremos a seguir).

Caso o usuário queira listar, de B10 até B200, quantas vezes aparecem números maiores que 38, a forma seria a seguinte:

=**CONT.SE(B10:B200; “>38”)**

Na CONT.SE, podemos usar alguns símbolos para criar as “expressões lógicas” como esta que vimos no exemplo acima... São eles:

< (menor que)

> (maior que)

<= (menor ou igual a)

>= (maior ou igual a)

= (igual a) – este não é necessário indicar.

<> (diferente de)

Note que o sinal de “=” (igual a) não é necessário ser colocado por um princípio muito óbvio:

A função =CONT.SE(B1:B10; “=100”) significa “Conte, de B1 até B10, quantas células tem conteúdo igual a 100. Ora, isso pode ser conseguido escrevendo, simplesmente, assim: =CONT.SE(B1:B10;100).

### 7.4.14.2. Funções de Soma Condicional

#### =**SOMASE(Int\_Cri;Critério;Int\_Soma)**

Essa função realiza uma soma condicional em que o usuário deverá informar segundo qual critério, em outro intervalo paralelo, deve ser encontrado para que se proceda com a soma dos

valores em um determinado intervalo. Veja:

|    | A        | B       | C         | D     | E     | F |
|----|----------|---------|-----------|-------|-------|---|
| 1  | Vendedor | Janeiro | Fevereiro | Março | Abril |   |
| 2  | João     | 40      | 5         | 60    | 60    |   |
| 3  | Ana      | 95      | 25        | 55    | 100   |   |
| 4  | Pedro    | 20      | 90        | 35    | 75    |   |
| 5  | João     | 90      | 30        | 40    | 40    |   |
| 6  | Mateus   | 70      | 15        | 65    | 20    |   |
| 7  | Mateus   | 25      | 90        | 95    | 80    |   |
| 8  | Ana      | 30      | 55        | 65    | 90    |   |
| 9  | Pedro    | 50      | 15        | 75    | 30    |   |
| 10 | Ana      | 70      | 50        | 30    | 75    |   |
| 11 |          |         |           |       |       |   |

Figura 7.45 – Uma planilha de controle de vendas.

Se o usuário quiser saber apenas quanto foi vendido, em Janeiro, por Pedro, basta informar =SOMASE(A2:A10;"Pedro";B2:B10). O Excel vai procurar, de A2 até A10, pela palavra Pedro, e, se encontrar, somará a célula equivalente da coluna B2 a B10.

Note que o primeiro intervalo é o que indica onde procurar o critério ("Pedro", no caso). Logo depois deve-se informar qual é o critério e, por fim, deve-se informar, depois do último ponto e vírgula, o intervalo da soma, ou seja, o intervalo onde estão as células que contêm os números a serem somados.

O critério da função SOMASE usa as mesmas regras do critério da função CONT.SE, e pode, também, utilizar expressões lógicas.

#### 7.4.14.3. Função SE

A função SE é a verdadeira "função condicional" do Excel. Por meio dela, é possível estabelecer dois valores de resposta possíveis e atribuir a responsabilidade ao Excel de escolher um deles.

Ou seja, você dá ao Excel: a faca, dois queijos e a razão que o fará escolher entre os dois!

A forma de usar a função SE é:

=SE(Condição;Valor Verdadeiro;Valor Falso)

onde

**Condição** é um teste, uma proposição a ser avaliada pelo Excel. Essa “proposição” só pode ter duas respostas: SIM ou NÃO. Note isso: o Excel não sabe qual é o estado desta proposição (ele irá avaliar no momento de acionar o ENTER).

**Valor verdadeiro** é a resposta que a função apresentará caso a condição seja verdadeira (caso a sua resposta tenha sido SIM).

**Valor falso** é a resposta que a função apresentará caso a condição seja falsa (ou seja, se sua resposta foi NÃO).

A condição do Excel sempre deve ser uma assertiva do tipo SIM/NÃO (ou booleana, como costumamos chamar); portanto, ela exige um operador de comparação entre dois valores. Um operador de comparação é um sinal usado para comparar dois valores. Os operadores que usamos são os mesmos das expressões lógicas das funções CONT.SE e SOMASE... Vamos lembrá-los:

< (menor que)

> (maior que)

<= (menor ou igual a)

>= (maior ou igual a)

= (igual a)

<> (diferente de)

Um exemplo muito comum na maioria dos cursos é a famosa planilha de notas dos alunos, que apresentará REPROVADO ou APROVADO de acordo com a média obtida pelo aluno. Veja um exemplo da planilha:

|   | A        | B      | C      | D     | E        |
|---|----------|--------|--------|-------|----------|
| 1 | Vendedor | Nota 1 | Nota 2 | Média | Situação |
| 2 | João     | 10     | 10     | 10    |          |
| 3 | Ana      | 9      | 9      | 9     |          |
| 4 | Pedro    | 8      | 7      | 7,5   |          |
| 5 | Mateus   | 8      | 9      | 8,5   |          |
| 6 | Paula    | 3      | 6      | 4,5   |          |
| 7 |          |        |        |       |          |

Figura 7.46 – Planilha de notas usando a função SE.

Observe que na coluna A estão os nomes dos alunos, nas colunas B e C estão as notas e na coluna D está a média (possivelmente calculada com o uso da função MÉDIA).

Na coluna E, deseja-se que o Excel apresente a palavra **Reprovado**, caso a média do aluno seja inferior a 7,0 (sete) e **Aprovado** caso a média do aluno seja igual ou superior a 7,0 (sete).

Para que o Excel faça isso, basta escrever a seguinte função SE (na célula E3, que corresponde ao campo “Situação” do primeiro aluno):

=SE(D3<7;"Reprovado";"Aprovado")

onde:

**D3<7** é a condição (também chamada teste lógico) que avalia se a média do aluno (localizada na célula D3) é menor que 7,0.

**“Reprovado”** é a resposta que a função apresentará caso a condição seja verdadeira;

**“Aprovado”** é a resposta da função caso a condição seja falsa;

Lembre-se: para usar textos dentro das funções do Excel, devemos escrevê-los entre aspas.

O resultado da função mostrada (após sua cópia com a alça de preenchimento) é:



|   | A        | B      | C      | D     | E         |  |
|---|----------|--------|--------|-------|-----------|--|
| 1 | Vendedor | Nota 1 | Nota 2 | Média | Situação  |  |
| 2 | João     | 10     | 10     | 10    | Aprovado  |  |
| 3 | Ana      | 9      | 9      | 9     | Aprovado  |  |
| 4 | Pedro    | 8      | 7      | 7,5   | Aprovado  |  |
| 5 | Mateus   | 8      | 9      | 8,5   | Aprovado  |  |
| 6 | Paula    | 3      | 6      | 4,5   | Reprovado |  |
| 7 |          |        |        |       |           |  |

Figura 7.47 – Resultado da função SE.

Quer entender definitivamente a função SE em português? Batize o primeiro ponto e vírgula de “*então*” e o segundo de “*senão*”, o resultado é que:

=SE(D3<7;”Reprovado”;”Aprovado”)

pode ser lida assim:

Se D3 for menor que 7, então apresente “Reprovado”, senão apresente “Aprovado”.

Fácil, não?

### 7.5. Construindo gráficos no Excel

Além das fórmulas e funções do Excel, é muito comum encontrar em concursos algumas perguntas acerca do recurso de criação de gráficos do Excel, que é muito simples de utilizar e entender.

Para construir um gráfico no Excel, é necessário selecionar uma sequência numérica em sua planilha e solicitar o comando de criação do gráfico.

|   | A        | B      | C      | D     | E        |
|---|----------|--------|--------|-------|----------|
| 1 | Vendedor | Nota 1 | Nota 2 | Média | Situação |
| 2 | João     | 10     | 10     | 10    |          |
| 3 | Ana      | 9      | 9      | 9     |          |
| 4 | Pedro    | 8      | 7      | 7,5   |          |
| 5 | Mateus   | 8      | 9      | 8,5   |          |
| 6 | Paula    | 3      | 6      | 4,5   |          |
| 7 |          |        |        |       |          |

Figura 7.48 – Basta selecionar o intervalo do qual se deseja construir o gráfico.

Depois de selecionadas as sequências textuais e numéricas que serão apresentadas no gráfico, basta acionar a ferramenta desejada do grupo **Gráfico**, na guia **Inserir**.

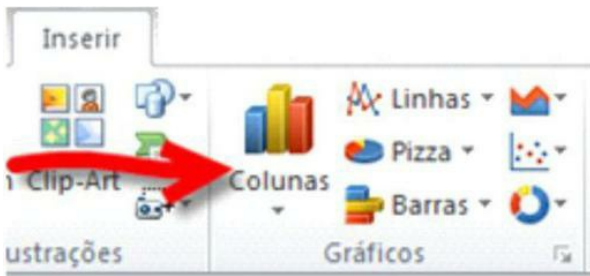


Figura 7.49 – Guia Inserir, grupo Gráfico.

Dependendo de qual ferramenta você escolha (para escolher o tipo do gráfico), surgirá um quadro de opções para escolher o subtipo dele. Escolhemos a ferramenta Colunas, e o primeiro

subtipo de colunas simples.

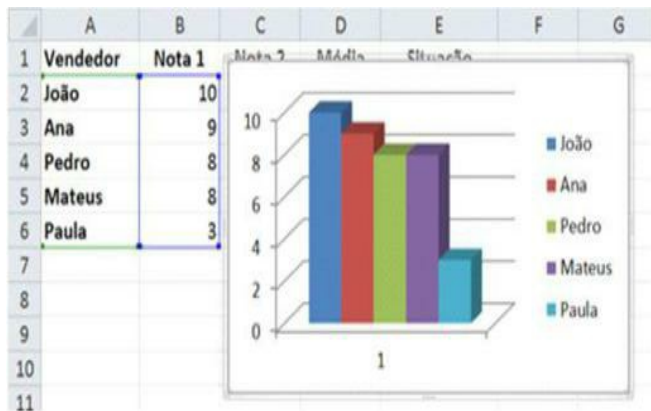


Figura 7.50 – Assistente de gráfico do Excel.

Quando o gráfico é selecionado, surge, na Faixa de Opções, um conjunto de guias próprio para Gráficos, contendo as guias **Design**, **Layout** e **Formatar**. Há algumas ferramentas interessantes em cada uma delas... Vá mexer!

Se quiser mais dicas acerca da criação e manipulação de gráficos no Excel, consulte o site da Editora Campus/Elsevier ([www.elsevier.com.br](http://www.elsevier.com.br)) na seção referente a este livro!



Figura 7.51 – Guias para Gráficos.

## 7.6. Outros comandos e recursos do Excel

Assim como no Word 2010, o Excel 2010 apresenta seus comandos em guias na Faixa de Opções. Vamos analisar alguns desses comandos (não todos, pois, em sua maioria, são semelhantes aos comandos de mesmo nome no Word).

Vou, contudo, apontar para os comandos que mais merecem nossa atenção em cada guia, ok? Considero que os demais, portanto, você já conhece do Word!

### 7.6.1. Guia Página Inicial

Não há muitos comandos no Excel que sejam diferentes em funcionalidade ou forma de acionamento que as ferramentas dessa guia no Word. Mas algumas, sim, pertencem somente ao Excel. Vejamos quais delas e onde:

#### 7.6.1.1. Grupo Alinhamento

No grupo Alinhamento, conforme se pode ver na figura a seguir, é possível encontrar as ferramentas de alinhamento normal de parágrafo (esquerda, centralizado e direita – note que não há justificado, como no Word). Também é possível ver os botões aumentar e diminuir recuo.

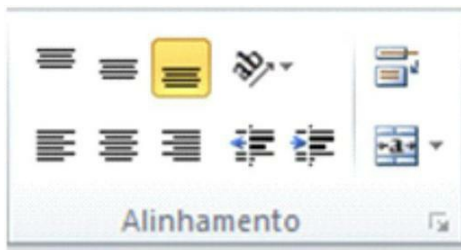


Figura 7.52 – Grupo Alinhamento.

Note, caro leitor, porém, a presença de algumas ferramentas interessantes:

- **Ferramentas de alinhamento vertical:** permitem alinhar o texto verticalmente na célula: em cima, no meio e embaixo (que é o normal).
- **Orientação (o “ab” inclinado):** permite escrever o texto de forma inclinada, ou vertical, ou uma letra em cima da outra... as várias opções se encontram quando se clica nesta ferramenta.
- **Quebrar texto automaticamente (o botão na extremidade superior direita, que parece um “robô com o braço esquerdo levantado e atirando raio laser pelos olhos” – tá, eu fui longe nessa, ein?):** permite escrever textos com mais de uma linha dentro de uma única célula.

- **Mesclar e centralizar** (o botão abaixo do robô, com a “letra a cercada por todos os lados”): Permite fundir várias células, transformando-as em apenas uma! Dentro do menu que se abre pela setinha à direita deste botão, há outras opções, como Dividir as células (desfazer a mesclagem).

Veja o resultado de algumas destas ferramentas:

|   | A       | B                            | C       | D     | E      |  |
|---|---------|------------------------------|---------|-------|--------|--|
| 1 |         |                              |         |       |        |  |
| 2 |         | Células Mescladas            |         |       |        |  |
| 3 |         |                              |         |       |        |  |
| 4 | em Cima | Célula<br>Quebrando<br>Texto | no Meio | Texto | Normal |  |
| 5 |         |                              |         |       |        |  |
| 6 |         |                              |         |       |        |  |

Figura 7.53 – Efeitos de Alinhamento.

Um clique no ícone de controle deste grupo (o pequeno ícone à direita do nome “Alinhamento”) abrirá a janela de **Formatar Células** diretamente na guia **Alinhamento** (guia, aqui, entenda-se: dentro da janela que se abrirá!).

Só um alerta: o ícone de controle do grupo **Fonte** também abrirá a janela de Formatar Células (com a diferença de que a guia aberta será a guia **Fonte**).

#### 7.6.1.2. Grupo Número

Oferece uma série de ferramentas para a formatação de números (estilos de números). Os números, no Excel, podem ser apresentados em diversos formatos: moeda, percentual, fração etc.

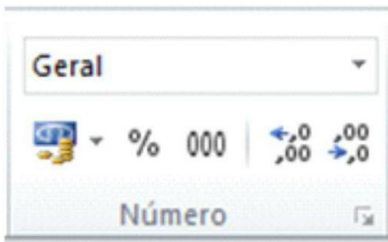


Figura 7.54 – Grupo Número.

Além do drop down (caixa de listagem), onde está aparecendo “Geral”, que permite escolher o tipo específico de estilo numérico (são mais de 10 estilos diferentes), é possível escolher os mais comuns nos botões abaixo:

- **Formato de Contabilização (o botão da cédula e das moedas):** apresenta o número com o seu símbolo de moeda (“R\$” no caso do Brasil!). Dá para escolher Dólar e Euro também (na setinha ao lado do botão).
- **Estilo de Porcentagem (o botão com o “% ”):** exibe o número como formato percentual.
- **Separador de Milhares (o botão do “000”):** formata qualquer número para ter um “.” (ponto) entre os milhares (de três em três dígitos) e duas casas decimais (separadas da parte inteira, claro, por uma vírgula).
- **Aumentar e Diminuir Casas Decimais:** a cada clique em cada um desses botões, fará o número de casas decimais apresentadas no número aumentar ou diminuir. Dica: as “setinhas” indicativas, nos desenhos dos botões, nos traem! Seta para a esquerda = Aumentar! Seta para a direita = Diminuir!

Veja o resultado:

|   | A                     | B         | C | D |  |
|---|-----------------------|-----------|---|---|--|
| 1 | Número                | 40        |   |   |  |
| 2 | Estilo Contábil       | R\$ 40,00 |   |   |  |
| 3 | Estilo Percentual     | 4000%     |   |   |  |
| 4 | Separador de Milhares | 40,00     |   |   |  |
| 5 | Aumentei Casas        | 40,0000   |   |   |  |
| 6 | Diminuí Casas         | 40,0      |   |   |  |
| 7 |                       |           |   |   |  |
| - |                       |           |   |   |  |

Figura 7.55 – Efeitos de formatação de números.

Perceba que o número 40 não se tornou 40%, mas 4000%! Claro! Os efeitos de formatação de número alteram o formato do número, não o seu valor! Se 40 virasse 40%, estaria mudando de valor (afinal, 40% é o mesmo que 0,4).

Na verdade, uma célula só muda de valor quando a gente (usuário) muda o valor da célula (escrevendo, manualmente, um novo valor).

E, para finalizar o grupo: um clique no ícone de controle deste grupo fará a abertura da janela **Formatar Células** (novamente), só que desta vez apontando para a guia **Número**.

#### 7.6.1.3. Grupo Estilo

Traz três ferramentas interessantes:

##### Formatação Condicional

Esta ferramenta permite que as células selecionadas apresentem algum efeito automaticamente de acordo com os valores delas.

É possível, por exemplo, especificar que uma determinada célula ficará em azul e negrito se o valor dela for maior ou igual a 7 e que ficará em vermelho e itálico se o seu valor for menor que 7 (note que isso – formatar condicionalmente – não é feito pela função SE, é feito por este recurso!).

Também é possível estabelecer indicadores (ícones ou barras coloridas) para apresentar, visualmente, a relação entre os valores de várias células: basta selecioná-las e escolher a opção certa dentro do comando formatação condicional.

Veja exemplos:

|   | A                                                                                | B  | C                                                                                 | D  | E                                                                                 |    |
|---|----------------------------------------------------------------------------------|----|-----------------------------------------------------------------------------------|----|-----------------------------------------------------------------------------------|----|
| 1 |                                                                                  |    |                                                                                   |    |                                                                                   |    |
| 2 |  | 10 |  | 15 |  | 35 |
| 3 |  | 15 |  | 10 |  | 30 |
| 4 |  | 5  |  | 30 |  | 20 |
| 5 |  | 30 |  | 25 |  | 10 |
| 6 |  | 25 |  | 20 |  | 25 |
| 7 |                                                                                  |    |                                                                                   |    |                                                                                   |    |
| 8 |                                                                                  |    |                                                                                   |    |                                                                                   |    |

*Figura 7.56 – Formatação condicional em ação.*

A coluna “A” mostrada acima usou o efeito de “Barras de Dados”, em que a formatação condicional cria barras coloridas crescentes de acordo com o valor das células selecionadas. As colunas “C” e “E” usaram o recurso de “Conjuntos de Ícones”, para representar símbolos diferentes de acordo com os valores das células.

### **Formatar como Tabela**

Transforma um conjunto (intervalo) de células em uma tabela. Esse recurso não só permite escolher a formatação da tabela (efeitos de cores e fontes para todas as células da tabela), como cria filtros para as colunas da tabela.



|   | A        | B       | C         | D     |
|---|----------|---------|-----------|-------|
| 1 | Colunas1 | Janeiro | Fevereiro | Março |
| 2 | NE       | 10      | 15        | 35    |
| 3 | N        | 15      | 10        | 30    |
| 4 | CO       | 5       | 30        | 20    |
| 5 | SE       | 30      | 25        | 10    |
| 6 | S        | 25      | 20        | 25    |
| 7 |          |         |           |       |

Figura 7.57 – Tabela formatada no Excel.

Preste atenção, também, ao fato de que uma nova guia é aberta quando você seleciona qualquer célula desta tabela: a guia **Design**, dentro de Ferramentas de Tabela.



Figura 7.58 – Guia Design, das Ferramentas de Tabela.

### Estilos de Célula

Ao clicar nesta ferramenta, surge uma galeria contendo formatos predeterminados de efeitos para células (você os aceita se quiser).

É um jeito rápido de aplicar efeitos às células. Além disso, depois de aplicar esses estilos, se quiser mudar na galeria a formatação, todas as células que têm esse estilo aplicado serão afetadas, apresentando o novo formato.

#### 7.6.1.4. Grupo Células

Traz as seguintes ferramentas:

- **Inserir:** permite inserir células, linhas, colunas ou planilhas inteiras.
- **Excluir:** consegue excluir células, linhas, colunas ou planilhas.

- **Formatar:** oferece diversos recursos de formatação da estrutura das células, linhas e colunas, como largura da coluna, altura da linha, ocultação ou exibição de células, linhas e colunas...

Esta ferramenta também oferece recursos para renomear e mover planilhas (em relação às demais planilhas da pasta de trabalho) além de proteger as planilhas com senha, para evitar inserção e apagamento de dados de forma não autorizada.

#### 7.6.1.5. Grupo Edição

Possui algumas ferramentas úteis, como o Localizar e Selecionar (que se assemelha ao Localizar do Word). Essa ferramenta apresenta apenas algumas opções a mais, como “Ir Para”, que permite navegar pela pasta de trabalho que se está utilizando.



Figura 7.59 – Grupo Edição.

Das ferramentas que se diferenciam do Word, podemos listar:

- **Classificar e Filtrar:** permite ordenar dados nas linhas do Excel com base no conteúdo de uma coluna (ordem alfabética ou numérica, sejam elas natural ou inversa) além de criar filtros (apresentar, na tela, apenas os dados que respeitem certo critério), o que é muito útil em planilhas grandes!
- **Soma (o símbolo do “Σ”):** dá acesso a um recurso que constrói rapidamente as funções (sem que seja necessário ao usuário digitá-las). Sinceramente, caro leitor, é uma ferramenta para “auxílio de preguiçoso”...
- **Preencher (o botão da “setinha para baixo”):** tem a mesma função da alça de preenchimento, ou seja: serve para ajudar a preencher a planilha facilmente.
- **Limpar (o botão da “borrachinha”):** permite limpar formatações, conteúdos, fórmulas e outros itens das células selecionadas.

#### 7.6.2. Demais guias do Excel

Vamos dar uma olhada em alguns dos recursos restantes do Excel, tendo em mente, claro, que

o mais importante sobre o programa já foi visto (é exatamente a parte que fala sobre cálculos).

As demais ferramentas fazem parte de diversas guias da Faixa de Opções do Excel:

### 7.6.2.1. Guia Inserir

#### Tabela Dinâmica

Este recurso, conseguido na guia *Inserir*, dentro do grupo *Tabelas*, serve para inserir uma tabela dinâmica na planilha (jura, é mesmo?).

Uma tabela dinâmica é uma forma de apresentar, resumidamente, dados que em uma planilha normal seriam considerados complicados ou difíceis de analisar. As tabelas dinâmicas oferecem uma maneira muito rápida de saber de informações precisas acerca do conteúdo apresentado na planilha.

Vamos à prática, analisando a imagem a seguir:

|    | A               | B              | C              | D                    | E |
|----|-----------------|----------------|----------------|----------------------|---|
| 1  | <b>Vendedor</b> | <b>Cidade</b>  | <b>Cliente</b> | <b>Valor Vendido</b> |   |
| 2  | João            | Santo André    | Be-a-Byte      | 960,00               |   |
| 3  | Ana             | Belo Horizonte | RedeGIR        | 1.350,00             |   |
| 4  | Pedro           | Florianópolis  | RedeGIR        | 580,00               |   |
| 5  | Mateus          | Belo Horizonte | RedeGIR        | 1.280,00             |   |
| 6  | João            | Belo Horizonte | Be-a-Byte      | 1.860,00             |   |
| 7  | João            | Belo Horizonte | Be-a-Byte      | 790,00               |   |
| 8  | João            | Florianópolis  | RedeGIR        | 600,00               |   |
| 9  | Ana             | Belo Horizonte | Be-a-Byte      | 930,00               |   |
| 10 | Ana             | Belo Horizonte | Eu Vou Passar  | 380,00               |   |
| 11 | João            | Santo André    | Be-a-Byte      | 1.140,00             |   |
| 12 | João            | Belo Horizonte | Be-a-Byte      | 1.680,00             |   |
| 13 | Mateus          | Belo Horizonte | RedeGIR        | 520,00               |   |
| 14 | João            | Belo Horizonte | Eu Vou Passar  | 1.800,00             |   |

Figura 7.60 – Planilha com dados crus.

Imagina se surgisse uma pergunta do tipo: “Quanto foi vendido, pela vendedora Ana, em Belo Horizonte?” ou, então “Quanto foi vendido ao Cliente ‘Eu Vou Passar’ de Santo André?”.

Tais perguntas são mais rapidamente respondidas por meio de uma tabela dinâmica, como a que está mostrada a seguir:

|    | A                  | B              | C             | D           | E           | F |
|----|--------------------|----------------|---------------|-------------|-------------|---|
| 1  |                    |                |               |             |             |   |
| 2  |                    |                |               |             |             |   |
| 3  | Soma de Valor Vend |                |               |             |             |   |
| 4  | Rótulos de Linha   | Belo Horizonte | Florianópolis | Santo André | Total Geral |   |
| 5  | Ana                | 14.780,00      | 3.970,00      | 5.280,00    | 24.030,00   |   |
| 6  | Be-a-Byte          | 6.130,00       | 1.690,00      | 2.000,00    | 9.820,00    |   |
| 7  | Eu Vou Passar      | 4.580,00       | 840,00        | 3.280,00    | 8.700,00    |   |
| 8  | RedeGIR            | 4.070,00       | 1.440,00      |             | 5.510,00    |   |
| 9  | João               | 15.550,00      | 4.380,00      | 9.090,00    | 29.020,00   |   |
| 10 | Mateus             | 14.140,00      | 9.270,00      | 5.910,00    | 29.320,00   |   |
| 11 | Pedro              | 19.810,00      | 6.280,00      | 7.050,00    | 33.140,00   |   |
| 12 | Be-a-Byte          | 8.630,00       | 2.600,00      | 790,00      | 12.020,00   |   |
| 13 | Eu Vou Passar      | 6.070,00       | 1.460,00      | 3.000,00    | 10.530,00   |   |
| 14 | RedeGIR            | 5.110,00       | 2.220,00      | 3.260,00    | 10.590,00   |   |
| 15 | Total Geral        | 64.280,00      | 23.900,00     | 27.330,00   | 115.510,00  |   |

Figura 7.61 – Tabela Dinâmica do exemplo anterior.

Exibe ou oculta todos os comentários das células da planilha. Quando esse comando é acionado, a barra de ferramentas Revisão é exibida.

Ao clicar em qualquer célula da tabela dinâmica, surgem as Ferramentas de Tabela Dinâmica, contendo duas guias: Design e Opções.



### 7.6.2.2. Guia Layout da Página

#### Ferramentas para Impressão

Na guia Layout da Página, encontram-se algumas ferramentas próprias para auxiliar o usuário no processo de impressão da planilha (processo este que, pelo fato de o Excel não ser WYSIWYG, é um pouco “melindroso”).

Só para lembrar: o Word é WYSIWYG (“What You See Is What You Get”, ou “O que você vê é o que você obtém”) porque aquilo que se vê na tela é exatamente o que se consegue impresso (afinal, no Word, temos, na tela, uma representação de uma página em branco, não é mesmo?).

No Excel, porém, não se imprime exatamente como se vê na tela, pois na tela temos uma planilha muito grande, que precisa de “certos ajustes” para caber numa folha de papel.

As ferramentas que vamos apresentar se encontram, todas, na guia Layout da Página, mas separadas em grupos diferentes... Preste atenção a elas:

#### *Área de Impressão*

Este comando, pertencente ao grupo **Configurar Página**, permite que você determine uma área específica que será impressa da planilha. Basta selecionar um conjunto de células e acionar o comando!

Depois de acionado o comando, uma borda tracejada irá circundar o intervalo de células determinado. Quando você acionar o comando Imprimir, perceberá que só aquela área específica será impressa (o restante da planilha será ignorado, mesmo que apresente conteúdo).

#### *Imprimir Títulos*

Este comando, também contido no grupo **Configurar Página**, dá acesso à janela de Configuração da Página, mais precisamente dentro da guia Layout, onde é possível, entre outras coisas, determinar quais linhas e colunas da planilha devem se repetir em todas as páginas quando a planilha for impressa.

Acompanhe um vídeo explicando a função deste recurso no hotsite deste livro, no site da Ed. Campus/Elsevier ([www.elsevier.com.br](http://www.elsevier.com.br)).

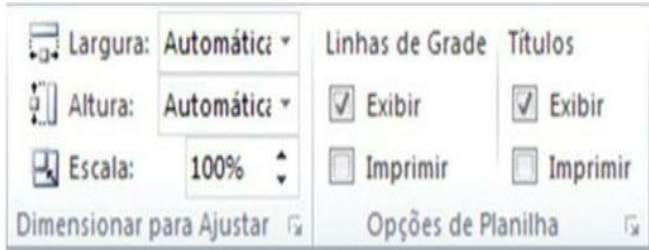
#### *Quebras*

Determina a posição das quebras de página (indicadores, na planilha, de onde se separam as páginas a serem impressas). Este comando também está localizado no grupo **Configurar Página**.

A planilha, caso seja grande demais, será impressa em várias páginas (folhas de papel) e, com isso, é necessário indicar, visualmente na tela, onde se encerra uma página e onde começa outra. Para alterar as posições destes indicadores, é necessário selecionar a célula na posição onde a quebra ocorrerá e inserir a quebra de página por meio deste comando.

#### *Dimensionar para Ajustar e Opções da Planilha*

Estes dois grupos estão juntos na guia Layout da Página e trazem algumas interessantes ferramentas.



*Figura 7.63 – Grupos Dimensionar e Opções da Planilha.*

No primeiro, podemos definir a largura (em páginas) e a altura (em páginas) que queremos para a nossa planilha ser impressa, ou determinar uma escala (em percentual) em relação ao tamanho atual da planilha.

Ao determinar uma largura e uma altura, a escala será automaticamente ajustada para caber melhor na quantidade de páginas definida.

No grupo Opções de Planilha, podemos solicitar se as linhas de grade (as finas linhas que separam as células) serão apenas mostradas na tela (opção “Exibir”) ou se serão impressas (opção “Imprimir” – o que, por sinal, não vai acontecer, pois está desmarcado).

O mesmo pode se dizer dos títulos das colunas (A, B, C, D etc.) e os títulos das linhas (1, 2, 3, 4 etc.). Normalmente, eles só são vistos na tela (opção “Exibir”), mas não são impressos (opção “Imprimir”, por isso ela está desmarcada).

### 7.6.2.3. Guia Fórmulas

Os comandos da guia Fórmulas dizem respeito, claro, à inserção de fórmulas e funções, além da análise e correção de erros que venham a acontecer nestas fórmulas.

A inserção das funções fica a cargo do grupo Biblioteca de Funções, mostrado abaixo, que separa as opções de funções de acordo com a categoria da função:



*Figura 7.64 – Grupo Biblioteca de Funções.*

Algumas ferramentas de facilitação, auditoria e análise das fórmulas podem ser encontradas nos demais grupos.



*Figura 7.65 – Mais ferramentas da guia Fórmulas.*

Através do grupo Nomes Definidos, é possível definir (atribuir) nomes amigáveis aos intervalos de células (como, por exemplo, dizer que as células A1:A20 serão chamadas de “Despesas” – desta forma, para somá-las, poderemos escrever **=SOMA(Despesas)**).

No grupo Auditoria de Fórmulas, podemos analisar como as fórmulas estão relacionadas, incluindo a apresentação de “setas” que rastreiam a relação entre a célula selecionada (que contém uma fórmula) e aquelas das quais ela depende (para as quais a célula aponta) – no comando Rastrear Precedentes.

Também é possível Rastrear Dependentes, que significa mostrar setas que apontem de uma célula para aquelas que contenham fórmulas que dependam dela.

|    | B              | C              | D                    | E | F          | G |
|----|----------------|----------------|----------------------|---|------------|---|
| 1  | <b>Cidade</b>  | <b>Cliente</b> | <b>Valor Vendido</b> |   |            |   |
| 2  | Santo André    | Be-a-Byte      | 960,00               |   |            |   |
| 3  | Belo Horizonte | RedeGIR        | 1.350,00             |   |            |   |
| 4  | Florianópolis  | RedeGIR        | 580,00               |   |            |   |
| 5  | Belo Horizonte | RedeGIR        | 1.280,00             |   |            |   |
| 6  | Belo Horizonte | Be-a-Byte      | 1.860,00             |   | 728.450,00 |   |
| 7  | Belo Horizonte | Be-a-Byte      | 790,00               |   |            |   |
| 8  | Florianópolis  | RedeGIR        | 600,00               |   |            |   |
| 9  | Belo Horizonte | Be-a-Byte      | 930,00               |   |            |   |
| 10 | Belo Horizonte | Eu Vou Passar  | 380,00               |   |            |   |

Figura 7.66 – Comando Rastrear Precedentes.

O comando Remover Setas faz com que todas elas sumam da tela (normalmente, ao solicitar que elas apareçam, em pouco tempo elas somem sozinhas).

O comando Mostrar Fórmulas faz com que as fórmulas escritas na planilha nunca sumam (nunca serão substituídas pelos resultados). Ou seja, mesmo depois do ENTER, as fórmulas continuam sendo mostradas. Basta outro clique nesta ferramenta para que volte ao normal.

#### 7.6.2.4. Guia Dados

Na guia Dados, encontramos vários comandos relacionados com o preenchimento das planilhas, bem como a análise e a correção dos dados nela inseridos.



Figura 7.67 – Guia Dados – primeiros grupos.



O grupo **Obter Dados Externos** oferece ferramentas para importar, de vários tipos de arquivos externos, os dados para a planilha, como arquivos do Access (banco de dados), arquivos de texto (do Word, por exemplo), arquivos da Internet (páginas), além de outras fontes, como arquivos XML, bancos de dados SQL entre outros.

O grupo **Conexões**, por sua vez, trabalha com os vínculos (links) entre a pasta de trabalho atual e as fontes de dados externas. Uma conexão é uma ligação entre o Excel e uma fonte de dados externa a ele, de modo que quando os dados foram alterados lá na origem, mantenham-se atualizados dentro da planilha do Excel.

Em resumo, o grupo Obter Dados Externos trabalha trazendo os dados para dentro da planilha (e tornando-os independentes do local de onde vieram). O grupo Conexões serve para “vincular” (ligar) o Excel ao local onde os dados estão, fazendo com que mantenham essa “ligação” (vínculo de dependência, mesmo) entre si.

Há, também, o grupo **Classificar e Filtrar**, com ferramentas mais específicas para o ordenamento dos itens de uma tabela e a apresentação seletiva (filtrada) dos dados desejados.

Interessante mesmo, porém, para uma prova, seria perguntar pelas ferramentas do grupo **Ferramentas de Dados**, que são várias e bem legais!



Figura 7.68 – Grupos Classificar e Filtrar e Ferramentas de Dados.

A ferramenta **Texto para Colunas** consegue dividir um texto que está em uma célula em várias colunas, desde que haja um símbolo no texto que possa ser substituído para tal objetivo (como um sinal de “,” por exemplo);

A ferramenta **Remover Duplicatas** serve para excluir as linhas da planilha que apresentam dados duplicados na coluna que for indicada.

O comando **Validação de Dados** oferece uma forma de determinar valores válidos para a inserção de dados nas células do Excel. Exemplo: pode-se determinar que, nas células da coluna “B”, só sejam aceitos valores numéricos entre 0 (zero) e 10 (dez).

O comando **Consolidar** serve para reunir os dados de vários intervalos (digamos A1:A10, B1:B10 e C1:C10) em um único intervalo de mesmo tamanho (digamos F1:F10) preenchendo-o com o resultado de uma função (soma, normalmente) de cada respectiva célula nos três intervalos.

Ou seja, a célula F1 terá a soma de A1, B1 e C1. A célula F2 terá a soma de A2, B2 e C2, e

assim por diante. Parece uma besteira (pois se poderia fazer tal coisa via uma função Soma na célula F1 e depois arrastando-a pela alça), mas é muito fácil de fazer e bem prático!

Ahhh! E tem uma diferença crucial: o Consolidar não mantém o vínculo (fórmulas), ele já traz como resultado do comando os números (ou seja, células preenchidas com valores numéricos, e não com fórmulas). Claro que você poderá escolher manter o vínculo ao realizar o comando, mas por padrão, ele não é mantido.

Dentro do comando **Teste de Hipóteses** há algumas opções, dentre as quais a mais interessante é **Atingir Meta**. Este comando permite que o Excel preencha uma determinada célula com um valor escolhido pelo programa, com o intuito de atingir um resultado determinado pelo usuário numa célula que contém uma fórmula.

Veja no exemplo abaixo:

|    | A                | B | C           | D | E        | F | G           | H |
|----|------------------|---|-------------|---|----------|---|-------------|---|
| 7  |                  |   |             |   |          |   |             |   |
| 8  | Valor do Carro   |   | Juros (Mês) |   | Parcelas |   | Valor Total |   |
| 9  | 38.000,00        |   | 2,60%       |   | 24       |   | 61.712,00   |   |
| 10 |                  |   |             |   |          |   |             |   |
| 11 | Valor da Parcela |   |             |   |          |   |             |   |
| 12 | 2.571,33         |   |             |   |          |   |             |   |
| 13 |                  |   |             |   |          |   |             |   |

Figura 7.69 – Exemplo de planilha para o Atingir Meta.

Explicando a planilha acima:

Nas células A9, C9 e E9 existem números (valores numéricos constantes);

A célula G9 tem a fórmula  $=A9*C9*E9+A9$  (tá, eu usei uma fórmula com juros simples, ok? Não reclama disso!);

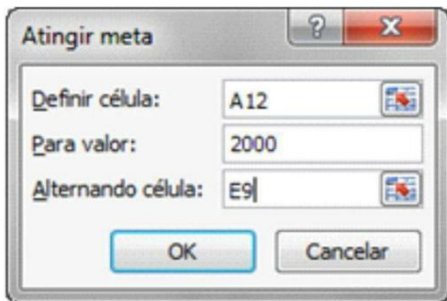
A célula A12 tem a fórmula  $=G9/E9$  (afinal, para calcular o valor da parcela, é necessário dividir o valor total do carro pelo número de parcelas).

Agora, imagine a seguinte situação: “Cara, ficou pesado para mim! Queria pagar mensalmente, no máximo, 2.000,00 reais!”. Sabendo que o nosso amigo chorão não se importa com o número de parcelas, nem com o valor total do automóvel, é só perguntar ao Excel: “Ei, Excel, que valor poderia haver em E9 (número de parcelas) para que o valor em A12 (valor de cada parcela) fosse 2.000,00?”.

É simples: a célula A12 tem uma fórmula que depende (diretamente) de E9 (número de parcelas). Então, para que A12 atinja a meta que queremos (2.000,00), alguém tem que “ceder”

(mudar de valor). Esse alguém indicado tem que ser uma célula contendo um número (não pode ser uma célula contendo uma fórmula).

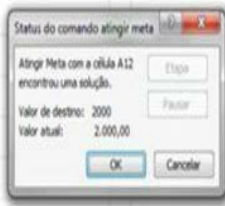
Veja como preenchemos a janela do Atingir Meta:



*Figura 7.70 – Atingir Meta sendo usado.*

Note, na figura abaixo, como a célula E9 mudou (para 37,549407) e isso fez com que o Valor Total do automóvel mudasse. Consequentemente, o valor da parcela mudou para o que se desejava (2.000,00) – na verdade, foi ele, o valor 2.000,00, quem fez todo o resto mudar!

|    | A                | B | C           | D | E         | F | G           | H |
|----|------------------|---|-------------|---|-----------|---|-------------|---|
| 7  |                  |   |             |   |           |   |             |   |
| 8  | Valor do Carro   |   | Juros (Mês) |   | Parcelas  |   | Valor Total |   |
| 9  | 38.000,00        |   | 2,60%       |   | 37,549407 |   | 75.098,81   |   |
| 10 |                  |   |             |   |           |   |             |   |
| 11 | Valor da Parcela |   |             |   |           |   |             |   |
| 12 | 2.000,00         |   |             |   |           |   |             |   |
| 13 |                  |   |             |   |           |   |             |   |
| 14 |                  |   |             |   |           |   |             |   |



*Figura 7.71 – Atingir Meta obteve uma solução.*

Há alguns outros comandos na guia Dados, mas estes são, sem dúvidas, os mais usados e, provavelmente, os mais cobrados em prova.

#### 7.6.2.5. Guia Revisão

Há vários comandos na guia Revisão, mas a maioria deles já foi vista no Word (os comandos são basicamente os mesmos), com exceção de poucos. Vamos, justamente, conhecer essas exceções!

#### Proteger Planilha e Proteger Pasta de Trabalho

Há dois comandos relacionados com proteção de dados.

Proteger Planilha impede que alterações sejam feitas nas planilhas, como inserção, modificação e exclusão de dados, alteração na largura das colunas e/ou na altura das linhas, inserção e exclusão de células, linhas e colunas etc.

Proteger Pasta de Trabalho impede as alterações na estrutura da pasta de trabalho, em si, como, por exemplo, inserção, movimentação e exclusão de planilhas.

Nos dois casos, é possível especificar uma senha que será usada para desproteger o item (desbloquear os limites impostos), permitindo, novamente, as alterações.

#### Compartilhar Pasta de Trabalho

Permite que uma pasta de trabalho (arquivo salvo pelo Excel) possa ser usada por mais de uma pessoa ao mesmo tempo (desde que tenha sido salvo, claro, num local em que todos tenham

acesso).

### Permitir que os Usuários Editem Intervalos

Depois de proteger uma planilha, é possível definir alguns intervalos de células que certos usuários poderão alterar (inserir, modificar ou excluir dados). Deste modo, você define *quem* pode alterar *onde* na sua planilha!

Esse recurso só funciona se a rede da qual o computador faz parte for uma rede corporativa Windows (um “domínio Windows”).

#### 7.6.2.6. Guia Exibição

A guia Exibição traz uma série de recursos interessantes, mas vamos, claro, nos ater aos que mais importam (os mais prováveis em prova).

O grupo **Modo de Exibição de Pasta de Trabalho** apresenta alguns comandos referentes à forma de apresentar a planilha na sua tela, começando pelo modo **Normal**, que é a forma como normalmente vemos a planilha.

Em seguida, temos o modo **Layout da Página**, que permite visualizar a planilha, na tela, conforme ela será impressa (em páginas, como no Word).

Há também, neste grupo, o botão **Tela Inteira**, que permite que a área da planilha passe a ocupar a tela inteira, sem a Faixa de Opções, nem a Barra de Fórmulas.

Há também o Grupo **Mostrar**, que permite exibir ou ocultar as linhas de grade, os títulos de linhas e colunas, a barra de fórmulas e a régua (visível somente no modo Layout da Página).

Há também o grupo **Zoom**, com comandos para aproximar e afastar a planilha, tornando seus conteúdos maiores ou menores respectivamente.

Finalmente, no grupo, **Janela**, vários são os comandos relacionados com a exibição das janelas onde o Excel se apresenta, como a possibilidade de criar novas janelas contendo a mesma planilha ou de organizar as janelas lado a lado.

Mas, no grupo Janela, o comando mais cobrado em prova é, simplesmente, o comando **Congelar Painéis**, em que é possível fixar uma ou mais linhas (na parte de cima) e/ou colunas (à esquerda) enquanto o restante da planilha mantém-se livre para a rolagem (arrastando pela barra de rolagem).

Basta clicar na primeira célula que deve ficar livre (exemplo, para fixar a linha 1 e a coluna A, seleciona-se a célula B2) e acionar o comando Congelar Painéis. Também há, dentro deste botão, os comandos **Congelar Primeira Linha** e **Congelar Primeira Coluna**, que facilitam o trabalho caso um desses seja o seu desejo.

Dentro deste botão também aparece o **Descongelar Painéis** quando algo já estiver congelado na planilha.

#### 7.7. Valores de erros (Mensagens #)

Algumas vezes, quando escrevemos uma fórmula no Excel, este não consegue dar um resultado correto e nos retorna mensagens precedidas pelo sinal de # (grade, sustenido, jogo da velha ou qualquer nome que queira dar). Essas mensagens são chamadas **Valores de erro**. Os

principais valores de erro e suas causas são listados a seguir:

- **#VALOR!:** é apresentado quando um usuário tenta inserir um argumento ou operando em uma fórmula que esta não entende. Exemplo, se na célula B3 existe 13 e na célula B4 existe “teste”, a fórmula =B3+B4 resultará em #VALOR!, porque na célula B4 existe um texto, que não pode ser calculado pela fórmula.
- **#DIV/0!:** ocorre quando a fórmula ou função tenta realizar uma divisão por 0 (zero).
- **#NOME?:** ocorre quando o Excel não reconhece o texto em uma fórmula. Por exemplo, quando se tenta inserir um nome de função que ele não conhece, como =ÇOMA(B2:B10). (Essa “doeu”, não foi?)
- **#REF!:** ocorre quando uma referência de célula não é válida. Como se quiséssemos que o Excel calculasse isto: =B2+A0. Esse erro ocorre quando usamos a alça de preenchimento e esta ultrapassa os limites das planilhas, tentando construir algo assim.
- **#NÚM!:** ocorre com valores numéricos inválidos em uma fórmula ou função. Por exemplo, quando uma função exige um argumento numérico positivo e o usuário fornece um argumento negativo.
- **#NULO!:** ocorre quando o usuário especifica uma interseção de duas áreas que não se interceptam. Em algumas funções e recursos do Excel, é necessário informar intervalos de células que se interceptam, caso o usuário informe intervalos que não possuem área de intersecção, esse erro será apresentado.
- **#####:** não é um erro na fórmula, mas sim um alerta de que o número apresentado na célula não cabe na largura da coluna. A solução é alterar a largura da coluna para que ele seja perfeitamente visualizado.

## 7.8. Referência circular

Esse é o tipo de erro causado pela falta de atenção do usuário. O Excel apresenta um erro de referência circular quando o usuário tenta inserir uma fórmula que dependa direta ou indiretamente da célula onde ela está sendo inserida. Veja o exemplo:

|   | A | B      | C | D |
|---|---|--------|---|---|
| 1 |   | 20     |   |   |
| 2 |   | =B1+B2 |   |   |
| 3 |   |        |   |   |
| 4 |   |        |   |   |

Figura 7.72 – Fórmula =B1+B2 escrita na célula B2.

O Excel tentará, em vão, resolver a equação solicitada pelo usuário, mas não será capaz porque isso ocasionará sucessivos cálculos fazendo o programa entrar em “loop”. O resultado é uma mensagem de erro e a célula apresentando 0 (zero) e uma seta apontando para onde o erro ocorreu.

Às vezes o problema não é tão “claro” de se ver, como poderemos notar no exemplo a seguir:

|   | A | B            | C | D |
|---|---|--------------|---|---|
| 1 |   | 20           |   |   |
| 2 |   | 30           |   |   |
| 3 |   | =SOMA(B1:B5) |   |   |
| 4 |   | 40           |   |   |
| 5 |   | 50           |   |   |
| 6 |   |              |   |   |
| - |   |              |   |   |

Figura 7.73 – A soma foi inserida em uma das células contidas no intervalo da qual é dependente.

Ainda é possível cometer o erro de referência circular apontando indiretamente para a célula em que a fórmula está inserida (mais difícil ainda de encontrar o erro).

|   | A      | B | C      | D |
|---|--------|---|--------|---|
| 1 | 10     |   | 20     |   |
| 2 |        |   |        |   |
| 3 | =A1+C3 |   | =A3+C1 |   |
| 4 |        |   |        |   |
| 5 |        |   |        |   |

*Figura 7.74 – A célula C3 aponta para A3, que, por sua vez, aponta de volta para C3 – referência circular indireta.*

## 7.9. Lembrando e aprimorando referências

No Excel, usamos as referências de células para indicar ao programa onde buscar dados, o que já vimos. Segue uma listagem mais completa das maneiras de fazer referências no programa.

### 7.9.1. Estilo de referência A1

Forma oficial de fazer referências no Excel. Neste estilo, as colunas são nomeadas por letras (de A até XFD) e as linhas são classificadas por números (de 1 até 10048576).

Já vimos essa forma de referência anteriormente, portanto, esta tabela serve como uma revisão:

| Para se referir a               | Use |
|---------------------------------|-----|
| A célula na coluna A e linha 10 | A10 |
| O intervalo de                  |     |



células na coluna  
A e linhas 10 a 20

A10:A20

O intervalo de  
células na linha 15  
e colunas B até E

B15:E15

Todas as células  
na linha 5

5:5

Todas as células  
nas linhas 5 a 10

5:10

Todas as células  
na coluna H

H:H

Todas as células  
nas colunas H a J

H:J

O intervalo de  
células nas colunas  
A a E e linhas 10 a  
20

A10:E20

Portanto, um usuário poderia solicitar a soma de todas as células da coluna B apenas digitando =SOMA(B:B). Não seria necessário fazer =SOMA(B1:B1048576).

### 7.9.2. Estilo de referência 3D

Este “estilo” nada mais é que uma implementação do estilo anterior para cálculos com dados em múltiplas planilhas.

O negócio é o seguinte: imagine um arquivo do Excel com seis planilhas (Plan1, Plan2, Plan3, Plan4, Plan5 e Plan6). Suponha que exista um valor em cada célula B10 dessas planilhas e que esses valores precisam ser somados, o que fazer?

Que tal assim?

= Plan1!B10 + Plan2!B10 + Plan3!B10 + Plan4!B10 + Plan5!B10 + Plan6!B10

Essa fórmula até que funciona, mas facilitaria muito se fizessemos assim:

=SOMA(Plan1:Plan6!B10)

Observe que o operador de intervalo (o sinal de dois-pontos) está entre os nomes das planilhas, o que constitui uma referência 3D. Pois é, por mais complexo que pareça, uma referência 3D é apenas um intervalo entre planilhas. Uma referência 3D inclui a referência de célula ou intervalo, precedida por um intervalo de nomes de planilhas, apenas.

### 7.10. Considerações finais

Bem, mesmo que você almeje um cargo público para a área jurídica e ache que o Excel não vai te ajudar em nada e não faz parte das atribuições do seu cargo, prepare-se... Todos os concursos, independentemente de cargo, salário, órgão ou nível de instrução, exigem Excel.

E, como pudemos perceber, ele não é uma ferramenta difícil. Se você conhecer esse programa bem, já sairá na frente dos seus concorrentes que não deram ao Excel a devida atenção.

## 7.11. Questões de Excel

1. No Excel, se o conteúdo  $=B1+C1$  da célula A1 for recortado e colado na célula A5, esta última normalmente deverá ficar com o conteúdo:
  - a)  $=B5+C5$ ;
  - b)  $=B1+B5$ ;
  - c)  $=C1+C5$ ;
  - d)  $=B1+C5$ ;
  - e)  $=B1+C1$ .
2. Na referência  $=[XXXX]YYYY!ZZZZ$ , a parte YYYY refere-se à:
  - a) célula onde a fórmula foi escrita;
  - b) planilha onde a célula foi escrita;
  - c) planilha onde está a célula ZZZZ;
  - d) planilha onde está a célula XXXX;
  - e) pasta de trabalho onde está a célula ZZZZ.
3. A fórmula  $=\$C1+B\$2$ , escrita em D5, quando copiada para G7, será reescrita como:
  - a)  $=\$C2+G\$4$ ;
  - b)  $=\$F3+E\$4$ ;
  - c)  $=\$C3+E\$4$ ;
  - d)  $=\$C3+E\$2$ ;
  - e)  $=\$F1+B\$4$ .
4. O resultado da função  $=SOMA(B2:B8)$  é:
  - a) o somatório dos valores localizados em todas as células entre B2 e B8;
  - b) o somatório dos valores contidos apenas nas células B2 e B8;
  - c) o maior valor contido nas células B2 e B8;
  - d) o maior valor contido em todas as células da coluna B;
  - e) o somatório dos valores de todas as células contidas na coluna B.
5. A fórmula  $=MÉDIA(B2:B6;C9)$  poderá ser escrita sem ocasionar erro de referência circular na célula:
  - a) B6;
  - b) B3;
  - c) C9;
  - d) C8;
  - e) B5.
6. Acerca das funções no Excel, julgue os itens a seguir.
  - I. A função MÁXIMO retorna o máximo divisor comum entre os argumentos dados.
  - II. A função MÉDIA calcula a média aritmética dos números dados no argumento.
  - III. Se houver alguma célula vazia na coluna B, a função  $=SOMA(B:B)$  a considerará 0

(zero).

Estão corretos apenas os itens:

- a) I;
- b) I e II;
- c) II e III;
- d) II;
- e) I, II e III.

7. Dadas as seguintes células de uma planilha Excel, com os respectivos conteúdos:

A1 = 1

A2 = 2

A3 = 3

A4 = 3

A5 = 2

A6 = 1

Selecione-se as células A1, A2 e A3 e arrastando-as simultaneamente, pela alça de preenchimento, sobre as células A4, A5 e A6, os conteúdos finais das células A1, A2, A3, A4, A5 e A6 serão, respectivamente:

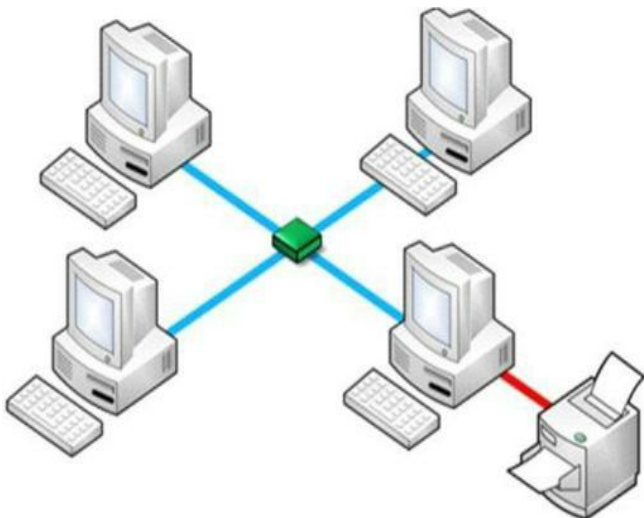
- a) 1, 2, 3, 1, 1 e 1;
- b) 1, 2, 3, 1, 2 e 3;
- c) 1, 2, 3, 3, 2 e 1;
- d) 1, 2, 3, 3, 3 e 3;
- e) 1, 2, 3, 4, 5 e 6.

### 8.1. Conceitos iniciais

A quantidade de informações que podem trafegar por um único computador é realmente imensa, imagine, então, quando são vários computadores reunidos... Uma **rede de computadores** é uma estrutura física e lógica que permite a conexão entre vários computadores com a finalidade de trocarem informações entre si.

Seguindo o conceito “bonito”, podemos dizer que “uma rede de computadores é um conjunto de módulos processadores (computadores), ligados por um sistema de comunicação, para permitir a troca de informações e o compartilhamento de recursos dos mais diversos fins”.

Para que haja uma rede de computadores, é necessário que existam, pelo menos, dois computadores e certos equipamentos capazes de conectá-los (fios, cabos, entre outros).



No exemplo da Figura 8.1, temos vários computadores interligados, e um deles está fisicamente conectado a uma impressora. Uma das vantagens da rede é que essa impressora poderá ser usada por todos os computadores dessa rede, em uma ação conhecida como **compartilhamento**. Compartilhar significa permitir que outros computadores usem um determinado recurso, como a impressora citada no exemplo anterior, que pertence, fisicamente, somente a um micro, mas poderá ser usada por todos os demais.

### 8.1.1. Classificação das redes

#### 8.1.1.1. Quanto à extensão

A maioria dos autores da área determina três classificações para as redes de computadores com relação à sua extensão (note bem que a diferença entre esses tipos de redes é meramente conceitual, não havendo uma unanimidade dos autores especializados ao apontar suas diferenças práticas), mas coloquei uma quarta (que, na verdade, aparece antes das outras):

**PAN (Personal Area Network – Rede Pessoal):** diz-se que uma PAN é uma rede em que todos os dispositivos envolvidos trabalham para **um único usuário**. É fácil imaginar isso quando nos lembramos daquele pessoal que carrega consigo diversos dispositivos eletrônicos como tablets, blackberry, celulares, máquinas fotográficas, headphones etc. Parecem até o Batman com um “cinto de utilidades” recheado de bugigangas.

• **LAN (Local Area Network – Rede Local):** uma rede de computadores de extensão pequena, normalmente dentro de um único prédio ou prédios vizinhos. Alguns autores afirmam que uma rede local se estende por, no máximo, 1 km.

• **MAN (Metropolitan Area Network – Rede Metropolitana):** uma rede de computadores em um espaço geográfico maior que o da LAN, mas ainda limitado. Ex.: rede de computadores no campus de uma universidade. Alguns autores definem o limite máximo de 10 km para uma MAN.

• **WAN (Wide Area Network – Rede Extensa ou Rede Geograficamente distribuída):** uma rede de computadores que não apresenta uma limitação geográfica. Exemplo: as redes de computadores dos grandes bancos e das operadoras de cartão de crédito, que se estendem pelo país todo, quando não pelo mundo!

#### 8.1.1.2. Quanto ao funcionamento

Essa classificação, em si, é só para fins didáticos, pois não serve para classificar a rede em si, mas a sua forma de trabalho (e isso depende exclusivamente da relação de interdependência entre os computadores – e programas – envolvidos):

• **P2P (Point-to-Point – Ponto a ponto):** uma rede na qual todos os computadores apresentam a mesma “importância” para o funcionamento da rede. Na verdade, é uma rede “cada um por si”, em que cada computador é responsável pelas informações que possui e deseja compartilhar com os demais. Ou seja, nessa rede, não se tem administração centralizada. Todos os micros ora “perguntam”, ora “respondem” – também chamada rede **homogênea**.

- **Client/Server (Cliente/Servidor):** nesta forma de funcionamento, define-se um (ou mais de um) computador para ser o centro das informações que se vão buscar. Esse computador (na verdade, mais precisamente um programa dentro desse computador) será chamado de **servidor** e deverá fornecer as informações aos computadores (novamente: programas) que as solicitarão (os **clientes**). Veremos mais acerca disso adiante.

Nota: é possível que cite esse assunto de ponto a ponto e cliente/servidor com o nome de Paradigmas de Funcionamento das Redes.

## 8.2. Sistemas de comunicação

A função de um sistema de comunicação é permitir a transmissão de dados entre dois componentes em uma rede, seja um sinal de telefonia, um arquivo de computador ou mesmo um programa de televisão. Vamos estudar agora os principais conceitos que envolvem o envio (transmissão) de sinais em um sistema de comunicação (rede).

### 8.2.1. Classificações da transmissão

Podemos classificar as transmissões de dados entre equipamentos por alguns critérios:

#### 8.2.1.1. Quanto ao tipo de transmissão

- **Análogica:** os sinais são transmitidos de forma analógica, ou seja, através de pulsos elétricos irregulares e contínuos, que podem assumir qualquer valor entre o mínimo e o máximo possíveis (é assim que são transmitidos, por exemplo, os sinais das linhas telefônicas convencionais).
- **Digital:** nesse modo de transmissão, os sinais são transferidos através de pulsos regulares (ou seja, com valores definidos) de energia elétrica. A diferença entre analógico e digital já foi mostrada com mais detalhes no início deste livro (na parte de hardware).

#### 8.2.1.2. Quanto ao sentido da transmissão

- **Simplex:** é uma transmissão que só acontece em um sentido (**de A para B**). Um exemplo seria a transmissão de TV, em que a emissora envia sinais e nossos aparelhos só conseguem captá-los (ou seja, a partir de nossos televisores, não podemos enviar dados para a emissora).
- **Half-Duplex:** a transmissão acontece nos dois sentidos (**de A para B e de B para A**), mas apenas em um sentido por vez. Ou seja, enquanto o “A” fala, o “B” não consegue falar, só escutar, e vice-versa. Um exemplo seria como funciona um walkie-talkie (ou o sistema de rádio da Nextel). Essa é a forma mais comum de transmissão nas redes locais de computadores.
- **Full-Duplex:** transmissão realizada nos dois sentidos simultaneamente. Os sinais podem trafegar, ao mesmo tempo, nos sentidos **de A para B e de B para A**. O melhor exemplo é o sistema telefônico.

### 8.2.1.3. Quanto à sincronização da transmissão

- **Síncrona:** a transmissão é sincronizada com o clock (frequência) da rede. Os dados são enviados de maneira sincronizada com o funcionamento da rede. Diversos dados são enviados sem intervalo entre eles, pois, como a transmissão está seguindo o “ritmo” da rede, o equipamento receptor será capaz de distinguir os diversos dados presentes na transmissão.
- **Assíncrona:** é uma transmissão que não está “moldada” segundo a frequência da rede, o que significa que ela não segue o “ritmo” da rede. Para esse tipo de transmissão, é necessário indicar onde começa e onde termina um caractere, um dado etc.

### 8.2.1.4. Quanto à comutação (chaveamento) utilizada na transmissão

O processo de comutação em um sistema de comunicação está relacionado à alocação de recursos para a transmissão. Essa “alocação de recursos” significa a escolha da rota (em alguns casos, rota física, mesmo), definição do tempo de uso dos recursos (ou seja, de quanto tempo aquele caminho será necessário) ou mesmo a questão da exclusividade do recurso (como o fato de as nossas linhas ficarem ocupadas, por exemplo, quando ligamos para alguém). Podemos citar, basicamente, três tipos de comutação para os sistemas de comunicação:

- **Comutação de Circuitos:** nesse tipo de comutação, antes de a transmissão iniciar-se, deve estar definido um caminho físico dedicado que ligue a origem e o destino. Ou seja, o sistema físico de telecomunicações envolvido deve estabelecer, primeiramente, a ligação física (circuitos chavearão para esse fim) entre o emissor e o receptor. É exatamente assim que procede a rede telefônica que nos atende atualmente: ou seja, quando você faz uma ligação telefônica, as redes da sua operadora (e da operadora do destino) “se viram como podem” para estabelecer as conexões físicas entre você e a pessoa para quem você está ligando.
- **Comutação de Pacotes:** sistema usado na Internet e em diversas tecnologias de redes de computadores. (Aliás, esse sistema é ideal para redes de comunicação de computadores, em que todos “querem poder” falar com todos.) Em uma rede de comutação de pacotes, não existem caminhos predeterminados para as mensagens a serem transmitidas – em vez disso, as mensagens são divididas em pequenas unidades de informação comumente chamadas de *pacotes*. Cada pacote leva consigo uma informação que diz de onde ele partiu e, mais importante, para onde vai. Cada nó (componente) do sistema lê essa informação do pacote (no seu cabeçalho) e o retransmite aos próximos nós até que chegue ao destino.
- **Comutação de Mensagens:** semelhante à comutação de pacotes (até mesmo porque a comutação de mensagens é antecessora da comutação de pacotes), não é necessário estabelecer previamente um caminho físico entre origem e destino para que a transmissão aconteça. A diferença dos pacotes é que a mensagem em si não é dividida. Na comutação de mensagens, um nó (componente da rede) só é “liberado” para a transmissão de outra mensagem quando finaliza a primeira, não importa o tamanho dela. Os pacotes, por sua vez, têm tamanhos limitados.

“João, essa parte da comutação ficou estranha... Não entendi...”

Calma, leitor. Mais adiante daremos atenção especial à comutação, especialmente à comutação de pacotes, que é importante para o entendimento do funcionamento das redes de



computadores atualmente (incluindo a Internet).

### 8.2.2. Problemas em uma transmissão

Nem tudo é um “mar de rosas” nos sistemas de comunicação e redes de computadores. Há diversos problemas, de ordem física, encontrados constantemente nesses cenários, que são uma verdadeira “pedra no sapato” dos administradores e gerentes desses sistemas.

A Esaf costumava exigir dos candidatos o conhecimento nesses defeitos, e, por causa disso, aqui vão eles:

- **Atenuação:** é uma consequência de a transmissão ser feita por meios físicos (fios, fibra óptica, ar etc.). A atenuação consiste na perda gradual da potência do sinal ao longo do meio de transmissão. Exemplo: quando gritamos, a “força” do nosso grito vai diminuindo à medida que o sinal sonoro se afasta de nós. Isso acontece também com a energia elétrica nos fios e com a luz nas fibras ópticas.
- **Ruído Térmico:** causado pela agitação dos elétrons em um condutor elétrico (fio). Esse tipo de ruído é constante em toda a extensão do condutor e é inevitável.
- **Ruído de Intermodulação:** causado pela presença de dois ou mais sinais de frequências diferentes em um mesmo condutor (um fio pode ser usado para transmitir diversos sinais diferentes em frequências variadas). Nesse tipo de ruído, uma transmissão em uma determinada frequência poderá induzir (e ser induzida) por um sinal transmitido em uma frequência próxima.
- **Ruído de Cross-Talk:** a famosa “linha cruzada” dos sistemas telefônicos. Esse ruído é causado pela indução eletromagnética que um condutor exerce sobre outro condutor próximo. Ou seja, vários fios dispostos lado a lado por uma longa extensão são mais suscetíveis a ruídos dessa natureza, pois um fio vai gerar um campo elétrico que irá induzir seus sinais em um condutor próximo (é exatamente como os fios das companhias telefônicas estão organizados).
- **Ruído Impulsivo:** é um ruído de grande amplitude (potência) que não é contínuo e surge sem previsão. Normalmente quanto há um distúrbio na rede elétrica, ou quando se liga um equipamento que consome grande potência (chuveiro elétrico, ar condicionado etc.), um pulso isolado de grande amplitude é gerado nos computadores (mais forte que o sinal que normalmente transita pela rede). É bastante difícil prevenir esse tipo de ruído. O ruído impulsivo não causa dano às transmissões analógicas (telefonia, por exemplo), mas é muito prejudicial às transmissões digitais (redes de computadores, por exemplo).

Lembre-se também de que a qualidade de transmissão de uma linha (um meio físico de transmissão, como um fio) é medida por uma razão entre a amplitude (força) do sinal e a amplitude do ruído (é a chamada razão sinal/ruído). Quando o ruído é muito alto (representando um percentual alto em relação ao sinal em si), a transmissão é classificada como de qualidade ruim.

### 8.3. Meios físicos de transmissão

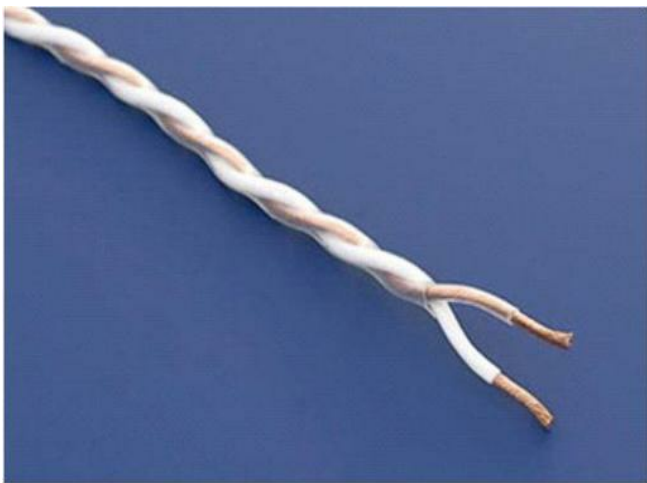
Para que haja transmissão de dados entre quaisquer dois componentes (computadores, por

exemplo), é necessário que haja meios por onde os sinais de dados (eletricidade, som, luz) possam passar. Esses meios de transmissão, que normalmente são cabos, serão apresentados agora:

1. Cabo par trançado;
2. Cabo coaxial;
3. Fibra óptica;
4. Ondas eletromagnéticas;

### 8.3.1. Cabo de par trançado

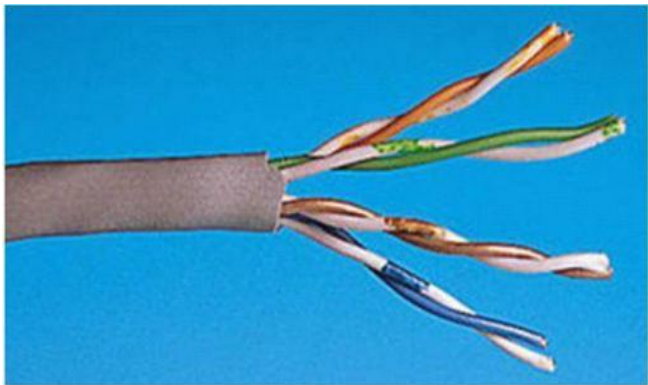
Conhecido também como simplesmente “par trançado” (twisted pair), esse cabo é amplamente usado em redes de comunicação de diversos tipos, tais como redes de computadores e redes telefônicas. Consiste em um (ou mais) par de fios trançados entre si (cada par tem seus dois fios dispostos como uma trança), para evitar o ruído de cross-talk.



*Figura 8.2 – Detalhe em um cabo par trançado.*

Os cabos atualmente usados não possuem necessariamente apenas um par, há cabos usados

em redes de computadores que usam até quatro pares de fios trançados.



*Figura 8.3 – Cabos de par trançado com quatro pares trançados (cada par é uma trança).*

Os cabos de par trançado podem ser classificados em dois tipos: UTP e STP.

#### **8.3.1.1. UTP – o cabo não blindado**

O cabo UTP (Unshielded Twisted Pair – ou “Par trançado não blindado”) apresenta-se como sendo a opção mais barata para os projetos da atualidade, e, por isso, a mais usada. Nesses cabos, as tranças não estão protegidas de interferências externas. A Figura 8.3 mostra um exemplo desse tipo de cabo. Ele é mais susceptível a ruídos externos, provenientes, por exemplo, de fontes eletromagnéticas fortes nas proximidades dos cabos.

Os cabos UTP são classificados por categorias, que indicam sua finalidade de uso (listei aqui apenas as mais comuns):

- **Categoria 1:** usado apenas em telefonia (são os cabos que chegam até nossos telefones partindo da companhia telefônica).
- **Categoria 5:** usado em redes de velocidades altas (100 Mbps) – como as atuais Ethernet –, mas suporta as redes de velocidades menores (10 Mbps).
- **Categoria 5e (5 enhanced – ou “melhorado”):** admite velocidades de transmissão muito maiores (até 1.000 Mbps) e é usado na terceira geração das redes Ethernet (chamada de Gigabit Ethernet).
- **Categorias 6 e 7:** usados em redes de velocidades de até 1.000 Mbps (Gigabit Ethernet).

“Ô, João, é necessário realmente decorar tudo isso?!” – você pode estar se perguntando. Ou melhor, me perguntando...

Sinceramente, caro leitor, não acredito que seja estritamente necessário saber tudo isso, mas se a Esaf ou a FGV (que é quem realmente tem “coragem” de exigir isso) vier a cobrar, especialmente nas suas provas mais “cabeludas”, como Auditor e Analista da Receita Federal e dos Fiscos Estaduais e Municipais, você já estará preparado.

Nem pense em decorar isto se você pretende fazer uma prova de tribunal ou feita pela FCC ou pelo Cespe/UnB (pelo menos, até agora, eles nunca fizeram uma questão tão técnica assim para quem não é da área de Informática).

“Certo, mas o que é Gigabit, o que é Ethernet?”

Calma... Veremos tudo isso! Gostaria de lembrar-lhe de algo que me acometeu agora: não se esqueça de que todos os cabos de par trançado (especialmente os que usamos atualmente) podem ser usados em segmentos máximos de 100 m (100 metros de cabo, apenas).

Dai a razão de os cabos UTP serem usados para apenas, normalmente, redes locais (LAN).

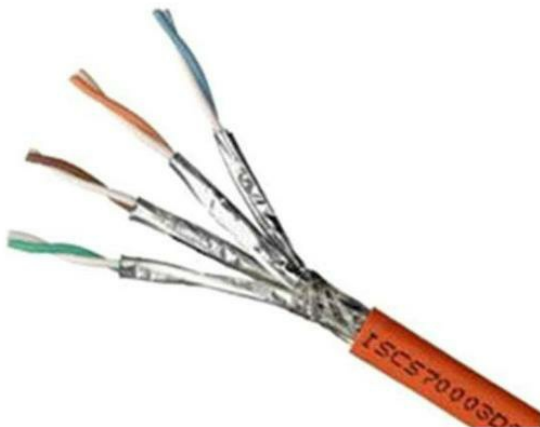
“Certo, mas se eu precisar ligar dois micros a uma distância de, digamos, 150 metros?”

É fácil, usamos um repetidor. (Calma... Direi o que é isso mais adiante, confie em mim!)

#### **8.3.1.2. STP – o cabo blindado**

O cabo STP (Shielded Twisted Pair – “Par trançado blindado”) é caracterizado por apresentar uma proteção (normalmente uma capa de material metálico – eu acho que é simplesmente “papel laminado”) que protege um par da indução de outros. Esse tipo de cabo é mais caro que o cabo UTP, e é menos flexível que este; portanto, em certos casos em que o “design” do projeto exige que o cabo seja bastante “dobrado”, o STP não será adequado.

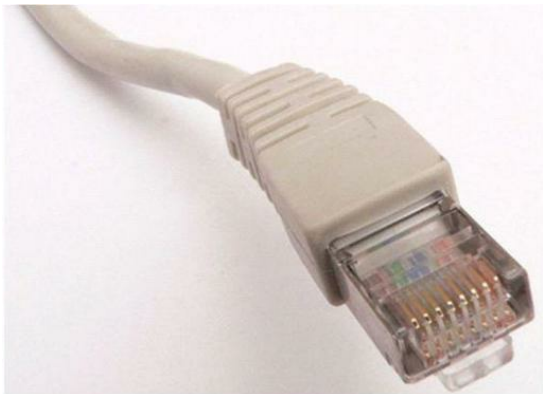
Sua proteção também garante mais imunidade a ruídos gerados por fontes externas, o que o torna recomendado para ambientes hostis, em que a emissão de ondas eletromagnéticas fortes é constante (fábricas, plataformas de petróleo, trios elétricos etc.).



*Figura 8.4 – Cabo STP – note a blindagem metálica.*

Tanto no caso dos UTP como nos STP, para que o cabo consiga “se conectar” a um equipamento qualquer, é necessária a presença de um conector (um pequeno dispositivo que faz a ligação dos fios presentes nos pares do cabo com o equipamento que se ligará à rede). Atualmente, o conector mais usado em redes de computadores é o RJ-45, feito de acrílico. Esse conector é bastante parecido com aquele conector usado nas linhas telefônicas (chamado RJ-11), mas é um pouco maior que este.

O conector RJ-45 é um pequeno cubo de acrílico com oito pinos metálicos em sua extremidade (onde as pontas dos fios do cabo UTP ou STP serão presas e com quem será realizado o contato elétrico para permitir a passagem dos sinais). Em resumo: cada um dos oito fios do cabo será conectado (por pressão) a um pino metálico localizado no conector RJ-45. E é através desses pinos (que farão contato com os fios) que a energia elétrica será conduzida de um componente da rede a outro pelo cabo.



*Figura 8.5 – Conector RJ-45.*

### 8.3.1.3 Os fios do cabo de par trançado

Vamos mergulhar dentro do cabo de par trançado (não importando, agora, se é o UTP ou o STP, pois ambos apresentam a mesma divisão estrutural dos fios). Sabemos que os cabos de par trançado são formados por oito fios internos trançados aos pares. Esses fios têm cores diferentes. (Não se prenda à ordem apresentada e sim às cores!)

Cada par de fios é formado por um fio de uma cor e outro fio da mesma cor mesclada ao branco. (Como um tracejado branco, um padrão de listras, uma cobra coral... como queira chamar.) Se a Figura 8.3 fosse colorida, você poderia facilmente identificar os quatro pares como sendo:

|                             |                       |
|-----------------------------|-----------------------|
| laranja/branco<br>& laranja | azul/branco &<br>azul |
| verde/branco &              | marrom/branco         |

verde

& marrom

“Ô, João, explica aí... É para decorar? E, além disso, para que servem as cores?”

Respondendo à sua primeira pergunta: na Esaf ou na FGV, talvez seja necessário um dia. Vá se preparando. A FCC, para a área geral (sem ser especificamente para Informática) parece não gostar muito do assunto. Cespe/UnB, por enquanto, nem sonhando!

E quanto à sua segunda pergunta. Para que servem as cores? Fácil! Já imaginou você montando um cabo de 30 metros de comprimento? Você arma a primeira ponta do cabo e conecta o primeiro RJ-45. Na hora de montar o segundo deles (a outra ponta) você vai querer arrumar os fios na mesma ordem que na primeira ponta. Como fazer isso?

Simples: olhando para a ordem dos fios em suas cores! Quero dizer que se você montou nesta ordem uma das pontas:

|                   |                     |
|-------------------|---------------------|
| 1. azul           | 5. laranja          |
| 2. branco & azul  | 6. branco & laranja |
| 3. verde          | 7. marrom           |
| 4. branco & verde | 8. branco & marrom  |

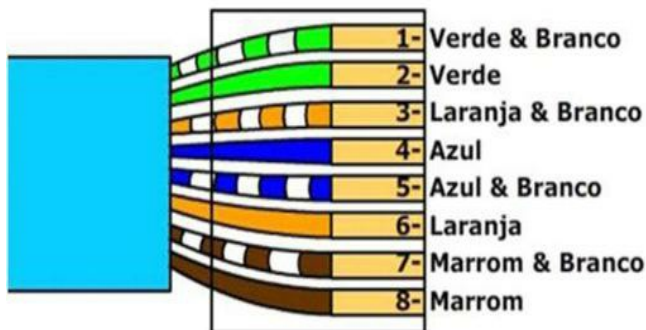
Basta repetir a dose na outra ponta do cabo para montar o chamado “cabo direto”, que é o mais usado em redes de computadores. A menos que você seja daltônico (feito eu), nobre leitor, é muito fácil localizar a ordem com que se montaram os fios.

“João, há alguma ordem para se seguir na montagem desses fios em relação aos pinos de um conector RJ-45?”

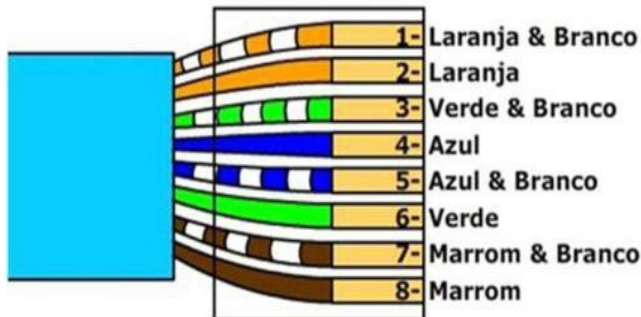
Sim, existem dois padrões de montagem, aprovados pela *EIA/TIA* (Electronic Industries Association/Telecommunications Industry Association) – Associação de Indústrias de Eletrônica/Associação de Indústrias de Telecomunicações. Recomenda-se que esses padrões sejam usados por todas as redes no mundo, para manter a compatibilidade entre as montagens dos cabos de rede – esses padrões são chamados de *EIA/TIA T568-A* (mais usado atualmente) e

**T568-B** (um pouco “esquecido” atualmente).

No 586-A, a ordem dos fios é mostrada na figura a seguir (o número que acompanha a cor do fio identifica o pino do conector RJ-45). A montagem 586-B é mostrada na Figura 8.7.



*Figura 8.6 – Configuração T568-A.*





Quero que você note que há apenas uma alteração entre os dois padrões: a posição dos pares verde e laranja. Observe que, na configuração A, o par verde está sendo destinado aos pinos 1 e 2 e, na configuração B, esse mesmo par está sendo usado para os pinos 3 e 6. O contrário, claro, acontece com o par laranja.

“João, explica uma coisa... Para que usar exatamente essas sequências? É para ser assim mesmo? Mas a eletricidade não ‘enxerga’ cores, não é?”

Precisamente! Não há necessidade de seguir essas ordens, basta que se saiba a relação entre os fios que ligam os pinos 1, 2, 3 e 6 (justamente os que se “alteram” de um padrão para o outro).

Deixe-me tentar explicar melhor: mesmo havendo quatro pares de fios (oito fios) num cabo de par trançado atual, só são usados (na maioria absoluta dos casos) quatro fios (dois pares). Adivinha quais? Os que estiverem ligados aos pinos 1, 2, 3 e 6!

Os pinos 1 e 2 são usados para a transmissão de dados (Tx, como alguns livros ainda chamam) e os pinos 3 e 6 são usados para a recepção de dados (Rx, segundo os mesmos livros). São dois pinos (e fios) para cada ação porque um dos pinos (e um dos fios) é usado para o positivo e o outro para o negativo (eletricidade, não é?). Essa é a regra usada pelos equipamentos que praticam comunicação na rede (as placas de rede dos computadores, por exemplo).

Ou seja, uma placa de rede de um computador vai sempre transmitir sinais pelos pinos 1 e 2 e sempre recebê-los pelos pinos 3 e 6.

Os fios ligados aos pinos 4, 5, 7 e 8 são “inúteis” para a montagem mais comum das redes de computadores. (Observe que existem montagens especiais, como as redes LAN em full-duplex, que usam os 8 fios da rede, mas elas são incomuns.)

#### 8.3.1.4. Cabo direto *versus* cabo cruzado

Devido ao fato de haver duas possibilidades de montagem dos fios em relação aos pinos (chamados de padrões de pinagem ou padrões de crimpagem), dois “modos de trabalho” para cabos de par trançado são possíveis:

- ***Straight Cable (Cabo Direto – Cabo Normal)***: esse cabo apresenta a mesma ordem dos fios internos em suas duas pontas (conectores RJ-45). Nesse tipo de cabo, o pino 1 de uma extremidade está ligado ao fio que levará exatamente ao pino 1 na outra extremidade, o mesmo para o pino 2 e os demais.

Ou seja, é fácil entender que este cabo tem, em seus dois conectores (nas duas pontas), padrões de pinagem iguais (ou 586-A em ambos ou 586-B em ambos). É usado para conectar um computador (placa de rede) a um Hub ou Switch (equipamento concentrador da rede).

Este tipo de cabo não pode ser usado para ligar diretamente dois micros (ou seja, duas placas de rede diretamente) porque, se a placa de rede do micro A transmitir sinais elétricos, eles vão sair pelos pinos 1 e 2 e também chegar, no micro B, nos pinos 1 e 2. (Isso é inadmissível, simplesmente porque a placa de rede do micro B também usa os pinos 1 e 2 para transmitir seus sinais.)

É por isso que existe o cabo cruzado.

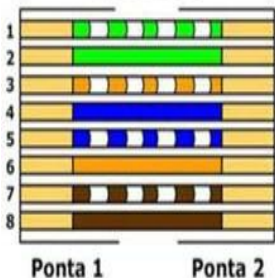
- ***Cross Over Cable (Cabo Cross Over – Cabo Cruzado)***: o pino 1 de uma extremidade está

ligado a um fio que vai levar ao pino 3 da outra extremidade, o pino 2 de uma ponta liga ao pino 6 da outra (resumindo: 1-3; 2-6).

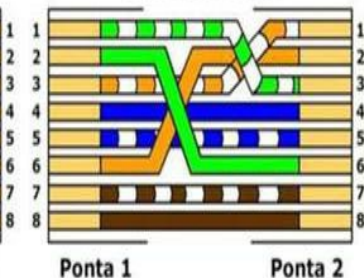
Isso faz um micro transmitir sinais por seus pinos 1 e 2 e esses sinais chegarem aos pinos de recepção do outro micro (3 e 6) – exatamente do que se precisa no caso de ligarmos os micros diretamente.

É perfeitamente possível notar, portanto, que a montagem desse tipo de cabo usa uma pinagem T568-A em uma das pontas e uma T568-B na outra ponta, não é mesmo?

**Cabo Direto**



**Cabo Cross Over**



*Figura 8.8 – Cabos UTP em suas duas montagens.*

### 8.3.2. Cabo coaxial

O cabo coaxial é formado por um condutor metálico central (que representa o polo positivo), envolto por uma malha metálica (polo negativo), que são, é claro, separados por um dielétrico (um isolante, como polietileno ou teflon).



*Figura 8.9 – Cabo coaxial (este é o cabo coaxial fino).*

Entre as características dos cabos coaxiais, podemos citar a sua baixa susceptibilidade a ruídos externos, sendo mais indicado que os cabos STP para ambientes “hostis” às comunicações. Há diversos tipos e medidas de cabos coaxiais usados em várias finalidades de comunicação. Havia praticamente dois tipos de cabos coaxiais usados em redes de computadores: o cabo fino (thin cable) e o cabo grosso (thick cable) – este último, muito antigo e sem uso atualmente.

Os cabos coaxiais são normalmente conectados a plugues (conectores) do tipo BNC, ainda usados hoje em equipamentos de vídeo profissionais (onde o cabo coaxial ainda é amplamente usado).



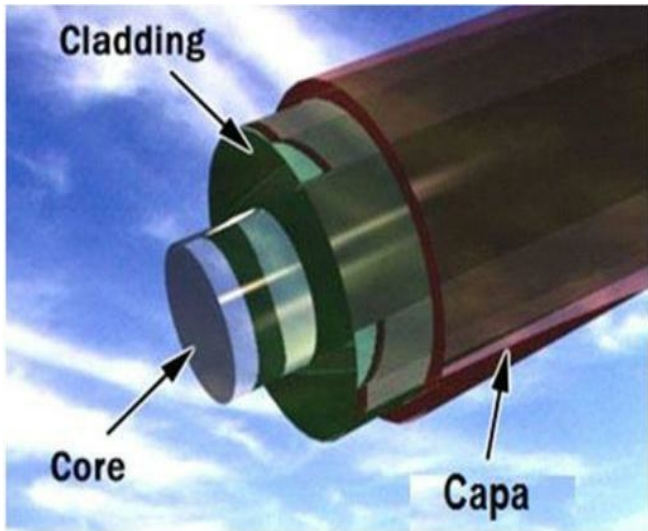
*Figura 8.10 – Conectores BNC.*

Os cabos coaxiais foram completamente substituídos pelos cabos de par trançado há, pelo menos, 15 anos (embora ainda haja empresas que teimam em manter estruturas com cabos coaxiais) – mas o que mais importa é que as bancas examinadoras ainda teimam em cobrar isso aqui! Francamente!

### **8.3.3. Fibra óptica**

Cabo usado para realizar a transmissão de pulsos luminosos (luz) em vez de sinais elétricos (como os cabos citados anteriormente). Ligado a uma extremidade de um cabo desses há um emissor de luz (que pode ser um LED – Diodo Emissor de Luz – ou um emissor de raio laser), à outra ponta do cabo, estará conectado um sensor, que detectará o sinal luminoso que transitou pela fibra.

O fio de fibra óptica é formado por um núcleo de vidro (o Core) por onde o sinal luminoso é transferido. Esse núcleo é envolto por uma camada de plástico que impede a passagem dos pulsos de luz (fazendo com que os raios reflitam sempre e não saiam do core). Essa camada é conhecida como bainha, ou casca (cladding). Externa à camada plástica, há a capa do fio, visível a todos nós.



*Figura 8.11 – Fibra óptica.*

Um cabo de fibra óptica apresenta, normalmente, um par de fibras (dois fios): um para transmitir os sinais em um sentido e o outro fio para transmitir sinais luminosos no sentido oposto (necessariamente, já que uma única fibra não poderá transmitir sinais nos dois sentidos). Mas, o mais comum, atualmente, é acumular vários fios de fibra óptica dentro de um mesmo cabo grosso, como mostrado na figura a seguir.



*Figura 8.12 – Cabo de fibra óptica (contém várias fibras dentro).*

As fibras ópticas podem ser basicamente divididas em fibras **monomodo** (single mode) e fibras **multimodo** (multi mode) – essa diferença se dá basicamente na espessura do núcleo (core) de vidro.

Uma fibra monomodo possui um core mais fino, que permite que a luz trafegue praticamente em linha reta. Sua principal característica é que a atenuação do sinal luminoso é menor, permitindo que haja mais comprimento útil de fio.

Uma fibra multimodo apresenta um core (núcleo) mais espesso, fazendo com que a luz “ricocheteie” nos limites do núcleo. São fibras mais baratas de fabricar e, conseqüentemente, de adquirir, mas o comprimento máximo do segmento deste tipo de fibra é bem menor que o da fibra monomodo.

Em uma fibra multimodo, vários raios de luz podem se propagar simultaneamente, pois poderão “resvalar” nas paredes do core em ângulos diferentes. As fibras multimodo experimentam transmissões semelhantes ao que se vê a seguir:

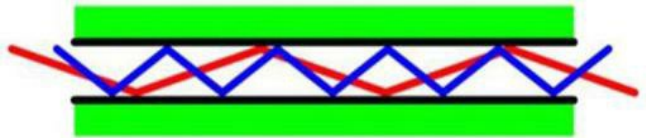


Figura 8.13 – Esquema de funcionamento da fibra multimodo.

“João, eu já ouvi falar em índice da fibra óptica, o que é isso?”

Existem basicamente dois tipos de fibras multimodo: a de índice degrau e a de índice gradual. Elas diferem entre si na forma como os feixes luminosos são refletidos de volta para o interior do núcleo da fibra.

A fibra óptica multimodo de *índice degrau* (step index) faz o feixe luminoso refletir de forma brusca, como um espelho faz diretamente a qualquer feixe que se aponta para ele. Isso acontece porque o índice de refração do núcleo da fibra é uniforme (homogêneo) em toda a sua extensão e bem superior ao índice de refração da casca (cladding).

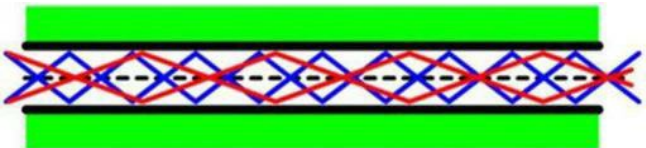
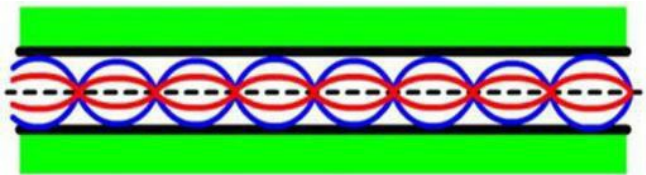


Figura 8.14 – Esquema de uma fibra óptica multimodo de índice degrau.

Na fibra de *índice gradual* (graded index), o índice de refração do núcleo da fibra cai à medida que vai afastando do centro do núcleo (gradualmente), causando um efeito de parábola no feixe luminoso (ou seja, o núcleo tem variados índices de refração).



*Figura 8.15 – Um exemplo simplificado da fibra multimodo índice gradual.*

Atualmente, as fibras de índice gradual são mais usadas no mercado, embora, teoricamente, sejam mais caras e complexas de fabricar que as suas irmãs de índice degrau.

Por fim, as fibras ópticas podem ser classificadas como monomodo. As fibras monomodo possuem um núcleo com espessura muito menor que a das fibras multimodo. Só para se ter uma ideia, a espessura do núcleo de uma fibra multimodo é de, normalmente, 50 a 200  $\mu\text{m}$  (micrômetros), enquanto a do núcleo da monomodo é de cerca de 3 a 8  $\mu\text{m}$ .

Como é muito estreito, o núcleo de uma fibra monomodo permite que a luz praticamente só viaje em um ângulo (em linha reta), atingindo o máximo em capacidade de transmissão e a distância que o feixe pode percorrer no cabo.



*Figura 8.16 – Um exemplo de fibra monomodo.*

#### 8.3.4. Ondas eletromagnéticas

Toda forma de transmissão que não utiliza fios para guiar os sinais entre emissor e receptor utiliza ondas eletromagnéticas como meio de transmissão.

De forma bem simplificada, as ondas eletromagnéticas são meios de transmissão que usam campos elétricos e magnéticos nos átomos (do ar e outras matérias) para a transmissão de sinais de vários tipos (como voz e dados). Essas ondas são, tipicamente, divididas por suas faixas de frequência em:

- **Ondas de Radiofrequência (ou RF):** são ondas eletromagnéticas com frequências situadas entre 30 MHz e 3 GHz. Nesse espectro de frequência, encontram-se várias tecnologias distintas como a transmissão de rádio e TV, as primeiras gerações de telefones celulares, entre outros.
- **Micro-ondas:** designam um espectro de ondas eletromagnéticas com frequências de 3 GHz a 30 GHz. A maioria das tecnologias atuais sem fio usa esse espectro de transmissão, como as tecnologias Wi-Fi, Bluetooth, Wi-Max e a telefonia celular atual.
- **Infravermelho (ou infrared):** espectro de frequências de ondas eletromagnéticas que se situa além da faixa das micro-ondas (ou seja, acima da frequência dos 30 GHz). Essa faixa de frequência se encontra no limiar da luz visível (um pouco abaixo da cor vermelha, a cor de frequência mais baixa que conseguimos enxergar).

Por serem luz (ou algo próximo ao que consideramos como tal), as ondas de infravermelho são obstruídas por corpos opacos, como qualquer objeto comum não translúcido. (Basta colocar a



mão na frente do controle remoto da TV para notar que ele não consegue controlá-la, porque simplesmente seus raios não chegam ao receptor.)

Isso demonstra a necessidade, para a transmissão em infravermelho, de *linha de visão* (ou *linha de visada*), que é a capacidade de o emissor e o receptor se “verem” sem a presença de qualquer obstáculo opaco entre eles.

“Ô João, é por isso que quando dois celulares, ou dois handhelds vão se comunicar por infravermelho, eles têm de ser posicionados um de frente para o outro?” – Precisamente! Eles têm de “olhar” um para o outro.

Hoje em dia, porém, caro leitor, não há mais tecnologias de transmissão que usem infravermelho além, claro, dos controles remotos! Celulares e outros dispositivos portáteis usam, em sua maioria, uma tecnologia de micro-ondas (Bluetooth).

Terminadas as principais “formas” físicas de comunicação, vamos analisar um pouco a teoria das topologias de redes (assunto mais relacionado às redes locais – LAN – apenas).

#### **8.4. Topologias de rede**

Não há forma mais fácil de explicar isto: topologia é um esquema, um layout, um formato que determina como os computadores vão se ligar entre si. Em uma rede LAN (pelo menos nas mais simples), normalmente escolhe-se uma única topologia (forma) para que os micros (também chamados de estações) fiquem ligados.

As topologias mais comuns são:

1. barramento (barra);
2. anel;
3. estrela.

##### **8.4.1. Topologia em barra (barramento)**

Em uma rede ligada em barra, todos os computadores estão ligados a um mesmo condutor central (um cabo, normalmente) compartilhado (ou seja, os micros usam o mesmo cabo, mas não simultaneamente).

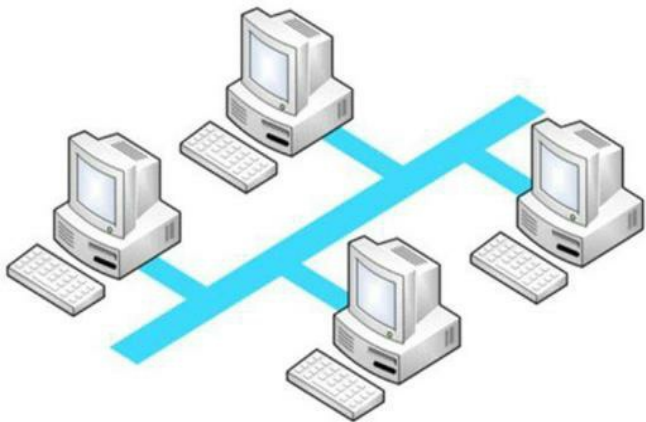


Figura 8.17 – Exemplo de rede barramento.

Devido à sua forma “limitante”, a topologia barramento apresenta algumas características interessantes, e muito fáceis de entender:

- **A rede funciona por difusão (broadcast):** ou seja, uma mensagem enviada por um computador acaba, eletricamente, chegando a todos os computadores da rede. Isso é ponto pacífico. O condutor central é um FIO! Um cabo! Ou seja, ele não tem condições de fazer outra coisa a não ser “mandar para todo mundo” os sinais elétricos que por ele trafegam.  
“É mesmo, João... Tem lógica!” – É claro que tem, amigo leitor! Qualquer sinal elétrico que um computador mandar para o condutor central vai chegar às placas de rede de todos os computadores ligados àquele condutor...

- **Baixo custo de implantação e manutenção:** devido aos equipamentos necessários (basicamente placas de rede e cabos). Essa característica é muito “relativa” porque hoje em dia, as redes barra, montadas fisicamente, não existem mais.  
As redes montadas fisicamente em barramento usavam cabos coaxiais, ou seja, só era possível criar redes realmente barra com cabos coaxiais. Como esse meio físico já está aposentado há uma longa data, não são mais vistas por aí redes barramento (pelo menos, não fisicamente).

- **Mesmo se uma das estações falhar, a rede continua funcionando normalmente:** pois os computadores (na verdade, as placas de rede, ou interfaces de rede) se comportam de

forma passiva, ou seja, o sinal elétrico é apenas recebido pela placa em cada computador, e não retransmitido por esta.

Também é fácil entender a razão dessa característica: o computador “A” envia algo através da rede barramento; a transmissão elétrica é enviada para todos (broadcast); o computador “B” estava desligado (opa!). Isso impede a mensagem de chegar aos demais, se estes estão ligados normalmente ao condutor central? Claro que não!

Não obstante o fato de todas as placas de rede (equipamentos presentes nos computadores que realizam a conexão física dos micros com os fios) conseguirem transmitir dados, elas são classificadas como passivas porque ao receber algum sinal pela rede, esse sinal não é retransmitido. Ou seja, ao receber qualquer sinal elétrico vindo de outro computador, uma placa de rede ligada em uma rede barra simplesmente analisa se aquele sinal é endereçado a ela. Se for, a mensagem é aceita (ou seja, o restante da mensagem é lido). Se a mensagem não for para aquela placa, ela simplesmente a descarta (não recebe o restante da mensagem).

“Ainda não entendi o porquê de não retransmitir, João!”

Simples, vamos a uma comparação básica: você está sentado numa sala de aula junto com seu colega, o **Fulano**. De repente, alguém entra na sala aos gritos: “QUEM É FULANO?! QUEM É FULANO?! VOCÊ PASSOU!!! SAIU O RESULTADO DO CONCURSO!!!”

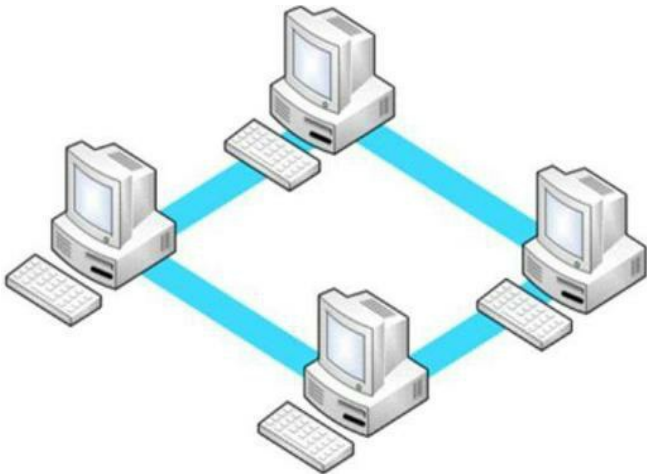
Bom... Mesmo que você tenha ouvido a mensagem e entendido que ela é para Fulano, é necessário falar de novo para ele? Ele estava do seu lado e o grito foi escutado pela sala inteira. Logo, você assumiu o comportamento passivo: escutou, viu que não era para você e descartou a mensagem. Não a retransmitiu!

Em suma, se uma mensagem é tipicamente em broadcast (difusão), uma placa tem plena consciência (se é que podemos usar esse termo) de que se a mensagem chegou até ela, é porque chegou a todas as demais também.

– ***Quanto mais computadores estiverem ligados à rede, pior será o desempenho (velocidade)*** da mesma (devido à grande quantidade de colisões). Para explicar melhor essa característica, vamos estudar mais adiante a ideia de ***colisão de pacotes*** em uma rede.

#### 8.4.2. Topologia em anel

Na topologia em anel, os computadores são ligados entre si em um caminho fechado (ou cíclico, como dizem alguns autores).



*Figura 8.18 – Exemplo de rede anel.*

Nesta topologia, as regras mudam bastante em relação à topologia barramento devido à própria forma como os sinais elétricos vão se propagar entre os micros. As principais características da topologia anel são:

- ***A mensagem enviada por um dos computadores atravessa todo o anel***, ou seja, quando um emissor envia um sinal, esse sinal passa por todos os computadores até o destinatário, que o copia e depois o reenvia, para que atravesse o restante do anel, em direção ao emissor.

“João, e por que voltar ao emissor?”

Para que ele saiba, quando receber o pacote enviado por ele mesmo, que a mensagem chegou a todos os micros da rede. Pois, se voltou a ele, atravessou todo o anel (todas as estações ligadas a ele).

“E isso muda alguma coisa com relação a um micro estar ‘desligado’ ou ‘quebrado’?”

Precisamente.

- ***Se um dos computadores falhar, toda a rede vai parar***: note que todo o anel é usado para a transmissão da mensagem em questão. E para que o computador emissor receba seu próprio pacote, ele deve passar (e ser retransmitido) por todos os computadores que formam aquele anel, dando às placas de rede desses computadores uma responsabilidade a mais:

receber; verificar se é para si; retransmitir.

• *Logo, se as placas de rede têm de retransmitir os sinais que recebem, elas apresentam um comportamento ativo.*

Analisando de forma semelhante ao broadcast. Você está lá, sentado ao lado do seu amigo: o Fulano. Alguém aparece à porta da sala de aula, mas, como o professor está no meio de sua explicação, o “cara da porta” não grita, mas sussurra ao ouvido do aluno mais próximo da porta: “Ei, você é o fulano? Se for, parabéns! Você passou. Se não for, avisa que ele passou!”.

O que vai acontecer depois disso? Simples: cada um que receber a mensagem vai repassá-la, sussurrando, ao vizinho, até que ela chegue ao Fulano. Se alguém se negar a passar a mensagem ou se tiver um súbito embaraço gástrico (vulgo “diarreja”) que o impeça de passá-la adiante, ela não chega. Eis aí a prova da necessidade de todos os nós (ou estações) envolvidos na transmissão do sinal em uma rede anel estarem funcionando.

As redes de topologia anel são, na realidade, montadas fisicamente de forma estrela. Ou seja, não houve, genericamente, redes anel realmente constituídas. As redes eram consideradas anel somente no funcionamento, não no esquema físico de ligação!

“Aí lascou tudo, João. Como é que uma topologia anel não é exatamente um esquema físico? Mas o significado de topologia não é de ser um esquema de conexão?”

Sim! Mas você há de convir, caro leitor, que a rede anel é muito “frágil” e “vulnerável”. Basta que um micro seja desligado para que a rede pare, não é mesmo? É interessante isso? Essa característica torna o uso das redes anel “convitativo”? Claro que não!

Então, as redes que funcionaram (e as que ainda funcionam) em anel são fisicamente estrela. Elas são apenas logicamente anel (aí surge a figura da “topologia física” e a “topologia lógica”). Ou seja, as tecnologias de rede que se dizem “anel” têm, na realidade, topologia física (ou seja, esquema físico, layout dos micros) em estrela, mas apresentam topologia lógica (funcionamento) em anel.

#### **8.4.3. Topologia em estrela**

Nesta topologia, os computadores estão ligados através de um equipamento concentrador dos cabos, o núcleo da rede, um equipamento que pode ser capaz de identificar o transmissor da mensagem de destiná-la diretamente para quem deve receber.

Se uma rede está realmente funcionando como estrela e se o equipamento central tiver capacidade para tanto, dois ou mais computadores podem transmitir seus sinais ao mesmo tempo (o que não acontece nas redes barra e anel).

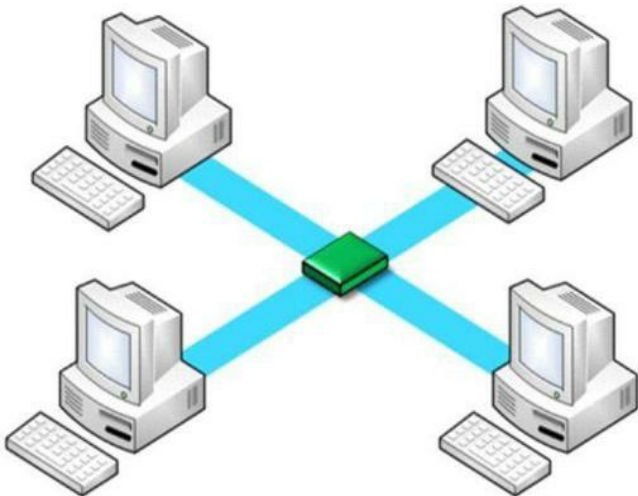


Figura 8.19 – Rede com topologia estrela.

As principais características a respeito da topologia em estrela que devemos conhecer são:

- **Admite trabalhar em difusão**, embora esse não seja seu modo cotidiano de trabalho.

Em uma rede estrela de verdade, é comum que um computador transmita um sinal (pacote) e este seja transmitido especificamente para quem deve recebê-lo. Seria, novamente, como o “cara” da porta que já sabe quem é Fulano. Ele bota a cabeça na porta, olha, olha, olha, acha o Fulano, vai até ele e, sussurrando só para ele, diz: “Cara, você passou! Parabéns!”.

Mas, eventualmente, as redes estrela podem trabalhar por difusão, especialmente quando o equipamento central (nó central, como é usado em geral) não souber quem é o destinatário (ou não tiver capacidade de ler a mensagem que está passando por si).

- **Todas as mensagens passam pelo nó central (concentrador)**. Essa, meu amigo leitor, é uma característica bem óbvia, tendo em vista a imagem anterior, que mostra como é uma rede estrela.

Mas vamos aprofundar isso: o concentrador (ou nó central) é um equipamento que recebe os cabos vindos de todos os computadores da rede e serve como um local para encaixá-lo, realizando, assim, a ligação física efetiva entre os micros.

Há basicamente dois equipamentos que assumem o papel de concentrador: o **hub** e o **switch**. Esses dois equipamentos são semelhantes fisicamente, mas bem distintos na forma como trabalham. Vamos ver isso mais adiante.

- ***Uma falha em uma estação (micro) não afeta a rede***, pois as interfaces (placas) de rede também funcionam de forma passiva.

“Ei, João, nessa não vou ficar calado. E quando as redes estrela estão funcionando como anel? (Você disse que isso era possível.) As interfaces de rede funcionam de forma passiva ou ativa?”

Essa é fácil! Se a rede funciona em anel (ou seja, apresenta topologia lógica anel), mesmo que seja fisicamente estrela, as placas de rede atuarão como na rede anel, ou seja, de forma ativa, ou seja, retransmitindo aos demais os sinais que receberem.

“Tá, mas isso torna os micros necessários, pois se um deles falhar, a rede para, não é?”

Não é bem assim. Sabendo dessa “vulnerabilidade” das redes anel, os fabricantes de componentes para essas redes colocaram no concentrador (nó central) dispositivos que detectam quando uma estação falha e, como consequência disso, fecham o anel sem a estação faltosa para que a rede continue funcionando.

Em suma, as redes anel são frágeis porque a falha em um micro faz a rede parar, mas como as redes anel sempre foram construídas fisicamente em estrela, essas “falhas” nunca pararam a rede, porque o concentrador (um hub específico para redes anel) sempre foi capaz de corrigir o problema, fechando o anel sem o micro falho.

“Quer dizer que os micros parados não param a rede? Quer dizer que a falha em um micro na rede anel não afeta os demais?”

Na prática, não afeta, porque as redes anel são fisicamente estrela – sempre! Mas ***se cair na prova***, não tenha dúvidas: ***uma falha em uma estação em uma rede anel vai parar a rede toda!*** Isso é uma característica das redes anel, sim!

- ***Uma falha no nó central faz a rede parar de funcionar***, o que, por sinal, é bastante óbvio. Aqui não tem o que discutir: se o concentrador (o equipamento central) falhar, a rede inteira vai falhar.
- ***Facilidade na implantação e manutenção***: é fácil ampliar, melhorar, instalar e detectar defeitos em uma rede fisicamente em estrela. Por isso essa topologia atualmente é a mais usada.

Atualmente, quando se fala em “essa rede é anel” ou “essa rede é barra”, na verdade, refere-se à topologia lógica, porque, em sua grande maioria, as redes atualmente são estrela física. E, na verdade, a topologia física que mais facilmente admite funcionamento em outros modos (ou seja, topologias lógicas) é a estrela.

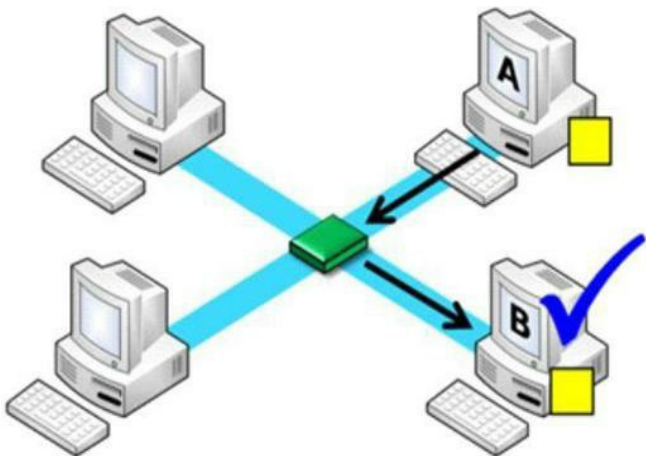
#### 8.4.4. Topologia física *versus* topologia lógica

Vamos analisar agora as variantes lógicas de uma rede estrela física. Começando, claro, com a própria rede estrela funcionando em estrela.

##### 8.4.4.1. Topologia lógica em estrela

Quando o equipamento central (o concentrador) é capaz de ler os sinais que trafegam por ele e

interpretar suas informações a ponto de saber direcioná-los para o destino específico, a rede física estrela funcionará como estrela lógica. É possível ver o envio de uma mensagem do micro “A” para o micro “B” na figura a seguir.



*Figura 8.20 – Rede estrela física trabalhando em estrela lógica.*

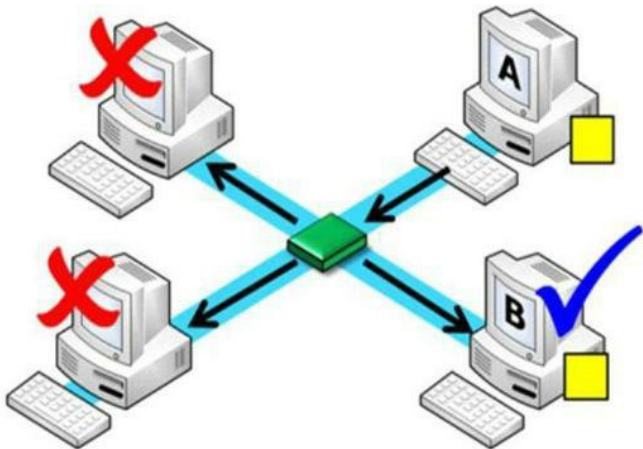
Essa montagem é possível quando o nó central é, por exemplo, um equipamento chamado **switch (comutador)**. Os switches têm a capacidade de ler os sinais (pacotes) que por ele trafegam e, com isso, enviá-los exatamente para o micro de destino.

#### **8.4.4.2. Topologia lógica em barramento**

Mas as redes estrela física também podem assumir outra configuração lógica, como barramento (a mais comum). Para tanto, basta que o equipamento central não saiba ler o sinal (pacote) que passa por ele. A mensagem impreterivelmente será retransmitida a todos os segmentos ligados àquele nó central (broadcast), já que ele não sabe filtrar nada.

Nesse caso, a mensagem chegará a todos os micros que, conseqüentemente, a rejeitarão (à exceção do micro de destino, que a aceitará). Nota-se o funcionamento exato de uma rede barramento.





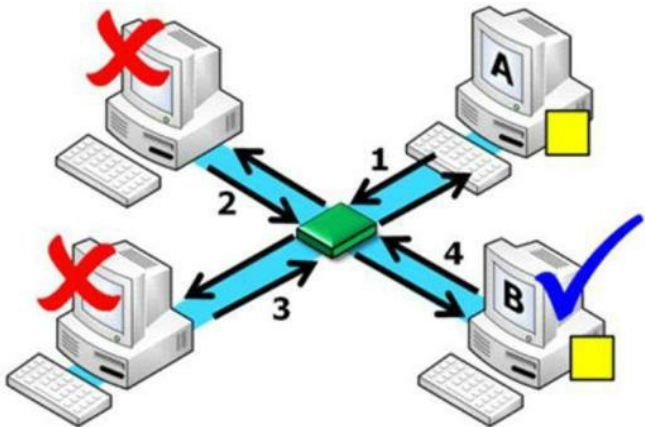
*Figura 8.21 – Rede estrela física versus barra lógica.*

O equipamento responsável por essa forma de trabalho chama-se **hub**. Um hub é um concentrador de cabos. Um hub não possui nenhum tipo de “filtro” ou “seletividade” para enviar os sinais aos micros que realmente devem recebê-los. Um hub simplesmente faz a cópia de todos os sinais que recebe e os envia na íntegra para todos os micros, portanto um hub funciona como aquele condutor central na rede barra física.

É simples assim: energia elétrica entra em uma das portas do hub e é replicada para todas as outras, como um T (um benjamin) desses de tomada elétrica.

#### 8.4.4.3. Topologia lógica em anel

Essa é mais rara hoje em dia. Mas quando havia redes em anel, elas funcionavam exatamente assim: fisicamente em estrela.



*Figura 8.22 – Rede estrela física versus anel lógica.*

É fácil entender a imagem anterior:

1. O micro “A” envia seu sinal na rede, objetivando o micro “B”; o sinal vai até o equipamento central específico para fazer o anel; este, por sua vez, envia o sinal ao próximo micro da sequência (a fim de dar continuidade ao anel).
2. O micro seguinte lê o pacote, vê que não é para si e o retransmite de volta ao nó central; este, novamente, envia ao próximo micro.
3. Esse terceiro micro lê o pacote, verifica que esse pacote não lhe pertence e o retransmite ao concentrador dessa rede; outra vez, o concentrador envia o pacote ao micro seguinte (que, no caso, já é o micro “B”).
4. O micro “B” recebe a mensagem e a lê, verificando que ela é mesmo direcionada a ele; o micro “B” a assimila (armazena e processa o pacote) e retransmite-o ao nó central para que dê continuidade ao processo de transmissão no anel; e este, para finalizar, envia o pacote de volta ao micro “A” (que é o próximo micro), fazendo, assim a transmissão ser encerrada.

### 8.5. Um pouco mais sobre comutação de pacotes

Como já foi visto superficialmente em um tópico anterior, a forma mais comum de comutação nos principais sistemas de comunicação de dados (como a Internet e as redes locais

de computadores) é a comutação de pacotes.

Comutação de pacotes é uma forma de comunicação que aloca (reserva) recursos de transmissão apenas para um pacote por vez, em vez de reservar recursos de transmissão para uma mensagem inteira (como acontece na comutação de mensagens) ou, pior, reservar a estrutura física real para dois computadores até que eles “decidam” parar de se comunicar (como seria o caso se fosse na comutação por circuitos).

“Tá, mas o que é um pacote?”

**Pacote é um “pedaço” da mensagem a ser transmitida.** Ou seja, antes de deixar o micro de origem, certos programas (chamados protocolos, como veremos adiante) dividem a mensagem em unidades menores, conhecidas como pacotes. (Datagramas ou quadros podem ser termos usados como sinônimos.)



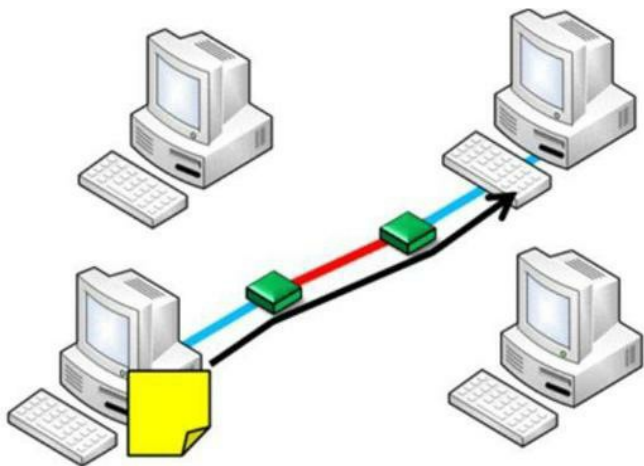
Figura 8.23 – Uma visão “poética” dos pacotes.

Como nessa figura, são consideradas mensagens todas as informações que podemos transferir por uma rede, como um e-mail, uma página Web (exemplo da figura), um arquivo, vídeos e músicas etc. No exemplo da figura, pode-se ver uma página Web (a página inicial do site [www.euvoupassar.com.br](http://www.euvoupassar.com.br)) em sua versão integral, e um “exemplo lúdico” da mesma página dividida em pacotes para que suas informações possam ser transmitidas pela Internet.

A comutação de circuitos contrasta, em muito, com a comutação de pacotes. Na comutação de circuitos, os componentes que irão travar a comunicação têm de, previamente, realizar uma

conexão física entre eles. (De algum modo, o sistema de telecomunicações que os liga tem de fechar circuitos, chavear linhas específicas, ajeitar conectores para que a energia elétrica saia do emissor e chegue ao receptor.)

Essa conexão física é dedicada (ou seja, nenhum outro micro vai usá-la em momento algum, somente os dois envolvidos). Veja um exemplo na figura a seguir.



*Figura 8.24 – Comutação de circuitos – um circuito físico dedicado entre os dois micros.*

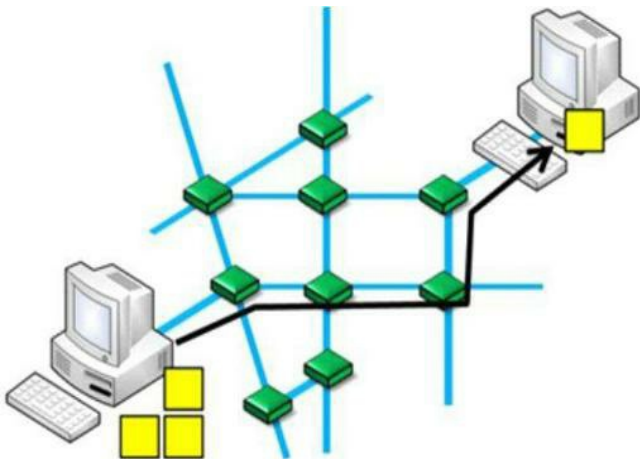
Felizmente, a comutação de circuitos não é usada em redes de computadores (nem seria adequada), porque um micro, ao mesmo tempo, pode estar “falando” com vários outros, e a comutação de circuitos impossibilita isso (razão pela qual, como já vimos, nossas linhas telefônicas ficam “ocupadas”). Como, nas linhas telefônicas, comumente falamos com apenas um interlocutor por vez, a comutação de circuitos é perfeitamente aceitável nelas.

“Ok, João, resume.”

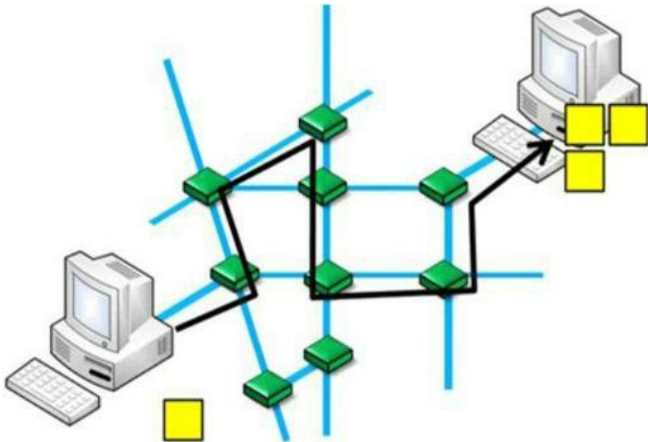
Claro! A comutação de circuitos exige um estabelecimento prévio, constante e ininterrupto de um circuito físico dedicado àquela transmissão. E até que a transmissão termine, os caminhos do emissor e do receptor estarão reservados a ela; portanto, inacessíveis a outras transmissões.

“Certo, e quanto às outras duas: pacotes x mensagens? O que dizer?”

Tanto na comutação de mensagens como na de pacotes, não há necessidade de estabelecimento de circuito físico dedicado entre a origem e o destino. As transmissões (sejam pacotes ou mensagens inteiras) trafegam por qualquer caminho que estiver livre na estrutura física de telecomunicações. (Essa estrutura não é dedicada, é compartilhada.) Os equipamentos intermediários (genericamente conhecidos como “nós de comutação”) vão estabelecer os caminhos da transmissão de acordo com regras específicas (protocolos de rota). Veja na figura a seguir.



*Figura 8.25 – Comutação de pacotes – um primeiro pacote sendo transmitido por uma rota.*



*Figura 8.26 – Comutação de pacotes – outro pacote, indo por outra rota.*

Mas a diferença principal entre as duas é o tamanho do que vão transmitir. Ou seja, o quanto cada sistema de comunicação vai se dedicar e alocar recursos para a transmissão. Senão, vejamos.

Como as mensagens têm tamanhos variados (algumas páginas têm poucos Kilobytes de tamanho, enquanto alguns arquivos já ultrapassam os Gigabytes), não dá para se “reservar” os recursos de transmissão do sistema (Internet, por exemplo) de acordo com o tamanho da mensagem – seria, basicamente, como reservar os espaços de um avião ou ônibus de acordo com o tamanho dos passageiros (gente maior em espaços maiores; gente menor em espaços menores – para mim, seria excelente: 1,90m; 130kg)!

Mas não é assim que acontece. Todas as pessoas (gente grande e pequena) têm de se “adequar” ao espaço determinado no avião (ou ônibus) – ou seja, todos têm de “caber” no tamanho da unidade de transmissão, um assento! O assento dos meios de transporte poderia ser entendido como um pacote. Se a informação é maior que o pacote, dois pacotes devem ser utilizados, mas o pacote não muda de tamanho para se ajustar à informação que tem de carregar.

Então, em uma rede que trabalha por comutação de mensagens, as mensagens maiores (como arquivos de vídeo sendo baixados) “travariam” (reservariam) para si a estrutura de comunicação por muito mais tempo que as mensagens simples (como um pequeno e-mail, por exemplo). E o

e-mail teria de esperar porque a rede em comutação de mensagens só se abre a uma nova transmissão depois de a transmissão da primeira mensagem ter sido realizada por completo.

“Ei, João, é como a fila do supermercado, não é? Em que uma pessoa quer passar apenas uma compra (um sabonete, por exemplo) e tem à sua frente um comprador com o carrinho inteiro descarregando no caixa.”

Exatamente! Eu não pensaria em uma comparação melhor, caro leitor!

Na comutação de pacotes, por sua vez, temos uma verdadeira democracia. Não importando se sua mensagem é maior ou menor que a minha. Quando um pacote da sua mensagem tiver sido transmitido, o sistema de comunicação (os nós de comutação da Internet, por exemplo) vai transmitir um dos pacotes da minha mensagem. Então, vai parecer que a minha mensagem e a sua mensagem são transmitidas simultaneamente pela estrutura da Internet (e são, só que de pacote em pacote).

Isso é possível porque todos os pacotes têm um tamanho máximo. São pequenos, em comparação às mensagens que carregam. Por exemplo, uma mensagem, como um arquivo de vídeo com, digamos, 1 GB, é dividida em cerca de 20.000 pacotes aproximadamente.

Seria mais ou menos como o supermercado do seu exemplo anterior, leitor. Só que o “cara do caixa” não atende de cliente em cliente, e sim de produto em produto. Ele passa um produto do cliente “A” (que tem muitos no carrinho); depois passa o sabonete do cliente “B” (que, com isso, vai embora, porque já passou tudo o que queria); depois volta ao cliente “A” para passar mais um produto (agora só faltam uns 149...); depois olha o cliente “C” com dois produtos na mão e passa um deles; depois volta pro cliente “A”... E assim por diante.

É democrático. Ninguém passa muito tempo esperando na fila, a não ser que tenha muito para passar. (É exatamente assim na Internet: arquivos maiores demoram mais para serem baixados.)

### 8.5.1. Agora, os pacotes em si!

Já sabemos que um pacote é um pedaço da informação a ser transmitida em uma rede que usa comutação de pacotes. Tudo bem! Mas eu gostaria, caro leitor, de ser um pouco mais específico e exato nesta descrição.

Todos sabemos que a manipulação e as transmissões de informação nos computadores são feitas de forma digital (através de pulsos que assumem apenas dois valores que representamos por 0 e 1). Ou seja, cada e-mail, cada página, cada arquivo que baixamos é apenas um conjunto de ZEROS e UNS.

Então, uma página, por mais bonita que seja, é, na verdade, um conjunto sequencial de 0 e 1, aliás, como qualquer outra informação. Sendo assim, podemos dizer que se:

01001010010101001010100001011110101010011101001001101001011010101000011101010  
é considerado uma mensagem (como um e-mail, por exemplo), então podemos concluir que:  
01001010010101001010100001011110101010011101001001101001011010101000011101010  
é um pacote dessa mensagem.

Claro que ninguém aqui vai ficar contando zeros e uns para entender o perfeito funcionamento das mensagens e dos pacotes, por isso vamos representá-los de uma forma mais “acessível” ao nosso entendimento.

Neque porro  
quisquam est  
qui dolore  
ipsum quia  
dolor sit amet,  
consectetur,  
adipisci velit...

| Origem                                            | Destino |
|---------------------------------------------------|---------|
| Neque porro<br>quisquam est<br>qui dolore<br>ipsu |         |

| Origem                                                         | Destino |
|----------------------------------------------------------------|---------|
| m quia dolor sit<br>amet,<br>consectetur,<br>adipisci velit... |         |

Mensagem

Pacotes

Figura 8.27 – Mensagem x Pacotes.

Note que os pacotes têm, juntos, todo o conteúdo da mensagem a que se referem, porque a função de dividir em pacotes é justamente para poder transmitir todos os dados da mensagem, mas em partes.

Observe também que há “coisa nova” nos pacotes, além do conteúdo da mensagem: é uma área que chamamos de **cabeçalho** (header) do pacote. Essa área é necessária, pois depois de criados os pacotes com o conteúdo dividido da mensagem, eles precisam ser identificados, de modo que consigam o seu objetivo: serem transmitidos com exatidão para o componente de destino correto. Portanto, no cabeçalho de um pacote, entre outras informações, estão o endereço do micro de origem e o endereço do micro de destino.

Tendo visto o necessário para entender pacotes, vamos analisar as principais tecnologias de redes de computadores atuais (e algumas não tão atuais). Afinal, é para isso que esse capítulo serve!

## 8.6. Arquiteturas de rede

Baseando-se nas três topologias vistas, várias empresas de tecnologia criaram seus próprios conceitos e definições a respeito de redes de computadores. A esses conjuntos de conceitos e características, damos o nome de **arquitetura de rede**.

Para que uma arquitetura de rede possa ser comercialmente usada, é necessário um processo de padronização por parte de algum órgão, instituto ou empresa desse gênero (como se passasse pelo selo do INMETRO para ser considerado seguro e pronto para o mercado). Na verdade, tudo relacionado à informática nasce em alguma empresa e deve passar pelo “crivo” da comunidade científico-comercial a fim de ser aceita como “usável”. IEEE, ISO, EITA, ITU são alguns dos órgãos que definem padrões mundialmente.

Quando um projetista de uma rede (a fim de montar a rede em sua casa ou empresa) define



que arquitetura vai utilizar, está definindo uma série de características sobre essa rede, por exemplo, desde o tipo de cabo utilizado até a topologia física e lógica da mesma (afinal, são características inerentes àquela arquitetura).

Em primeiro lugar, vamos analisar algumas arquiteturas utilizadas (atualmente e antigamente) em redes locais (LANs).

### 8.6.1. Ethernet (IEEE 802.3)

A arquitetura de rede conhecida como Ethernet, definida pelo padrão 802.3 do IEEE (Instituto de Engenheiros Elétricos e Eletrônicos) é, sem dúvida, a mais utilizada atualmente. Consiste em ligar computadores em uma topologia de barramento (lógica), permitindo, assim, o acesso de todos eles ao meio de transmissão. (Lembre-se de que “barramento” é um caminho necessariamente compartilhado.)

As redes Ethernet já foram montadas fisicamente em barramento, ou seja, com cabos coaxiais e conectores BNC, mas, atualmente, é mais comum encontrar essas redes montadas fisicamente em estrela, fazendo uso de cabos de par trançado e hubs ou switches. Em suma, uma rede Ethernet pode apresentar sua topologia física como barramento ou estrela, mas sua topologia lógica (funcionamento) será sempre barramento.

Também é possível encontrar variações da Ethernet com fibra óptica, o que traz a possibilidade de aumento da distância entre as estações envolvidas.

As redes no padrão Ethernet originalmente (pelos idos da década de 1980 até o início da década de 1990) se conectavam a uma velocidade de 10 Mbps (megabits por segundo) e hoje já permitem taxas de transmissão bem superiores. As redes Ethernet de segunda geração (também conhecidas como **Fast Ethernet**) transmitem dados a uma taxa de 100 Mbps. O padrão mais novo de Ethernet transmite dados a 1.000 Mbps (o equivalente a 1 Gbps – gigabit por segundo), por isso é conhecido como **Gigabit Ethernet**.

Existe uma forma para determinar as características de uma rede Ethernet usando apenas uma sigla. Na verdade, essa sigla define um padrão, regulamentado pelos órgãos competentes na área de comunicação de dados. Eu chamo simplesmente de **VbaseC** (onde V é Velocidade e C é o tipo do cabo usado na rede).

- **10Base2**: uma rede no padrão Ethernet montada com cabo coaxial fino e que usa a velocidade de 10 Mbps (a distância máxima entre uma estação e outra é de 185 metros). Por usar cabo coaxial, a topologia física desse padrão é barramento. Ele é bastante antigo e **não é mais usado**.
- **10Base5**: uma rede que usa cabo coaxial grosso e velocidade de 10 Mbps (a distância máxima entre uma estação e outra, nesse tipo de cabo, é de 500 metros). Uma rede nesse padrão também usa topologia física de barramento. **Esse é o padrão Ethernet mais antigo de todos**.
- **10BaseT**: uma rede de 10 Mbps que usa cabos de par trançado categoria 3 ou superior (**T é justamente de trançado**). A distância máxima entre a estação e o hub é de 100 metros (limite do cabo). Por usar cabos UTP, a topologia física desta rede é estrela (utiliza hub ou switch como nó central).

- **10BaseF:** uma definição que especifica qualquer rede Ethernet de 10 Mbps que utiliza fibra óptica como meio de transmissão (duas fibras – uma para transmitir, outra para receber). Há vários subpadrões com diferenças sutis entre eles (10BaseFX, 10BaseFB, 10BaseFP). A distância entre as estações é uma das características que variam de acordo com esses subpadrões. A topologia física dos padrões 10BaseF é estrela.
- **100BaseTX:** uma rede Fast Ethernet (100 Mbps) que usa cabos de par trançado categoria 5. Nesse padrão, o cabo UTP usa apenas dois dos quatro pares. A distância máxima entre a estação e o Hub é de 100 metros (limitação do cabo). Apresenta topologia física em estrela. Esse padrão é muito utilizado atualmente, com hubs (ou switches) como nós centrais da rede.
- **100BaseFX:** uma rede Fast Ethernet (100 Mbps) que usa dois cabos fibra óptica (um para transmitir e um para receber). A distância máxima entre as estações é de 2.000 metros. A topologia física deste padrão Ethernet é estrela.
- **1000BaseT:** uma rede Gigabit Ethernet (1.000 Mbps, que é o equivalente a 1 Gbps) que utiliza cabos de par trançado UTP categoria 5, 5e ou 6. Por usarem cabos que já são amplamente difundidos em redes Fast Ethernet, a “migração” para esse padrão de 1.000 Mbps é mais fácil (a maioria das redes de computadores montadas atualmente já é neste formato). A distância máxima entre estação e hub é de 100 metros (que é o limite do cabo). A topologia física deste padrão, claro, é estrela!

#### 8.6.1.1. Como funciona a arquitetura Ethernet?

Por funcionar em barramento, a rede Ethernet transmite seus sinais por difusão; portanto, uma estação (computador) emite seus sinais na rede (no meio) e esses sinais elétricos são enviados a todas as outras estações. Embora todas as estações recebam o sinal elétrico, somente o computador que tem o endereço correspondente ao endereço de destino do pacote irá aceitá-lo. Isso está descrito de forma perfeita na **Figura 8.21** (na qual é descrito o funcionamento das redes estrela física e lógica barramento).

A esse tipo transmissão, com já vimos, chamamos de difusão (Broadcast) e ela acontece simplesmente porque o hub não consegue ler as mensagens que passam por ele (o hub Ethernet, na verdade, é apenas um “barramento” para ligar todos os micros); portanto, ele não sabe quem é o destinatário das mensagens.

É mais simples do que parece: a placa de rede do computador emissor (aquele que vai enviar seus sinais) joga o sinal elétrico no meio de transmissão (cabos e hub), e o hub, por sua vez, faz chegar o sinal a todos os computadores ligados a ele. O **quadro** (nome dado ao **pacote** aqui) é lido por todas as placas de rede dos computadores da rede, mas somente uma o aceitará (porque seu endereço coincide com o endereço apresentado no quadro que foi transmitido). Os demais computadores simplesmente descartam o quadro porque ele não estava endereçado para eles.

Esse funcionamento também abre possibilidade de uma forma de interceptação dos dados bem interessante: uma placa de rede pode funcionar em “modo promíscuo”, ou seja, ela pode capturar todos os quadros que a ela chegam, independentemente do endereço contido no mesmo. Isso faz um computador receber dados que não eram direcionados a ele. Costumo chamar o “modo promíscuo” de “ouvido de fofoqueira” – ou seja, seu micro será capaz de captar todos os

pacotes que a ele chegam, mesmo que não sejam a ele direcionados oficialmente.

“Ei João, esse modo promíscuo serve para quê?”

Veremos isso mais adiante, no capítulo sobre segurança da informação.

Ok, Ok! O que está demonstrado na Figura 8.21 é a forma como os computadores reagem a um pacote que foi transmitido por alguma outra estação da rede. Mas como “essa estação” obteve o direito de “falar” na rede, ou seja, como ela conseguiu lançar seu pacote na rede?

### 8.6.1.2. CSMA/CD

Em todas as arquiteturas de redes de computadores, há regras a serem seguidas pelos computadores para estabelecerem conexão uns com os outros. Essas regras são genericamente conhecidas como *protocolos*. Isso seria como as “boas maneiras” da rede, ou seja, as regras de convivência que os computadores seguem.

Existem diversos protocolos para vários objetivos específicos, mas vamos nos ater a um tipo de protocolo: o protocolo de acesso ao meio (ou método de acesso ao meio), que define exatamente como os computadores vão ter acesso ao meio físico para “falar” através dele.

“Ô, João, dá para fazer uma daquelas tuas comparações?”

Claro! Imagine que, em uma sala de aula, um professor está realizando uma explicação sobre um tema muito complexo. Em um dado momento, um aluno tem uma dúvida em um trecho da explicação. O que ele faz para “acessar o meio físico” e fazer sua pergunta? O que ele faz para usar o ar da sala para transmitir seus sinais sonoros aos tímpanos do professor?

Por questão de educação (boas maneiras, ou “protocolo”), ele *levanta uma mão* esperando que o professor lhe dê a chance da palavra.

Note que podemos fazer uma série de considerações acerca desse exemplo:

1. O professor falando para todos está transmitindo em broadcast.
2. O aluno não pode fazer a pergunta simultaneamente à explicação do professor, porque eles não se entenderiam. (Os sons iriam se misturar no ar e todos receberiam uma informação audível, mas incompreensível.)
3. Para ter acesso ao meio (o ar) o aluno tem de ser o único a transmitir, e ele só terá certeza disso se o professor lhe autorizar. Razão pela qual ele “pediu permissão” – ou “protocolou um pedido para falar”.
4. Quando o aluno fizer sua pergunta em voz alta ao professor, estará transmitindo também em broadcast, pois todos os alunos o ouvirão.

Trazendo isso novamente para o mundo das redes de computadores (e, mais precisamente, as redes Ethernet), temos que:

1. O computador que estiver transmitindo, tem que fazê-lo sozinho, em broadcast.
2. Se um computador quiser transmitir um pacote, deve ter certeza de que o fará sozinho (ou seja, tem de se certificar de que nenhum outro computador esteja usando atualmente o meio físico – os cabos).
3. Quando dois ou mais computadores tentam transmitir sinais (isso é possível, mas não desejável), os sinais elétricos se misturam, causando um efeito desagradável e fazendo os dois micros entenderem que a “tentativa de transmissão” foi malsucedida – isso é o que chamamos de *colisão*.

Como a Ethernet é uma rede que baseia seu funcionamento em barramento, a transmissão acontece necessariamente em broadcast (isso é fato!). Como o broadcast vai acontecer de todo jeito (não dá para evitar), é necessário que os protocolos que regem a rede Ethernet saibam lidar com o broadcast. E eles sabem!

Na arquitetura Ethernet, o controle sobre como as informações serão enviadas entre os computadores é regido pelo protocolo **CSMA/CD**. Na verdade, todas as características de funcionamento das redes Ethernet se baseiam no conceito de CSMA/CD.

CSMA/CD significa Carrier Sense Multiple Access with Collision Detection, ou “Sensor da Portadora com Acesso Múltiplo e Detecção de Colisão”. Todas as placas de rede Ethernet trazem em sua memória ROM, desde a fábrica, esse conjunto de regras armazenado. Ou seja, todas as placas de rede Ethernet vão obedecer aos “preceitos” descritos nesse protocolo de acesso.

Vamos fazer uma análise mais simples desse protocolo, apenas separando em duplas as letras de sua sigla. Assim fica bem mais fácil, eu garanto!

- **CS (Carrier Sense – Escutar a Portadora):** os computadores que desejam transmitir dados ficam “escutando” o movimento da rede, quando esta se “acalma” é momento de jogar suas mensagens no meio.

A chamada “portadora” é apenas uma frequência que transita pelo meio (cabos da rede) sem conduzir mensagens. A portadora é a “correnteza” que permite aos computadores enviar seus dados. Quando os computadores detectam a presença da portadora no meio, é sinal de que eles podem transmitir porque não há mais nada sendo transmitido, ou seja, o meio de transmissão está livre.

Em relação ao exemplo da sala de aula, é como se o aluno que deseja fazer a pergunta ficasse esperando o momento certo para transmitir – esse momento se dá quando o professor se cala (para beber água, por exemplo) ou quando ele dá a deixa, perguntando, por exemplo, “alguma dúvida?”.

Ou seja, o aluno (ou micro) fica lá, morrendo de vontade de transmitir, mas espera pacientemente a sua chance, que seria a deixa do professor ou a presença da portadora pronta e liberada só para ele.

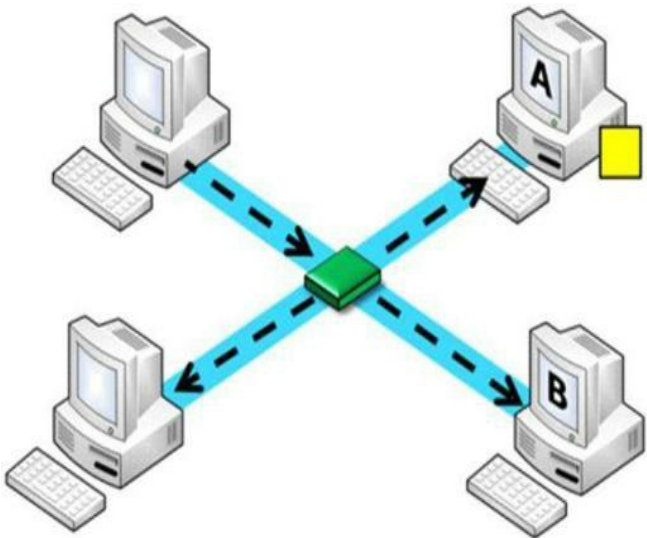


Figura 8.28 – Micro “A” querendo transmitir (mas não pode, pois a rede ainda está sendo usada).

Quando a rede for liberada, ou seja, quando nenhuma transmissão mais estiver acontecendo no meio físico, a placa de rede do micro que deseja transmitir vai sentir isso (sentir a portadora) e vai lançar seus sinais no meio. O resto da história, você já conhece (broadcast).

Mas ficar apenas “escutando” para ver se é possível transmitir tem um “efeito colateral”: quem disse que o aluno do exemplo será o único a fazer uma pergunta quando o professor der a chance?

- **MA (Multiple Access – Acesso Múltiplo):** significa que vários computadores podem tentar o acesso ao meio, basta que a portadora (frequência para transmitir sinais) seja detectada por eles.

É o caso de haver mais de um aluno querendo fazer uma pergunta ou de vários micros estarem prontos para transmitir quando a portadora for liberada. Isso realmente não é legal!

“João, o que acontece quando dois (ou mais) micros tentam acesso ao meio físico ao mesmo tempo?”

Eles misturam seus sinais elétricos no meio, formando uma transmissão estranha, deturpada, que não será entendida. (Como duas pessoas fazendo perguntas diferentes ao professor simultaneamente – duvido que ele entenda.) Essa ocorrência é chamada de *colisão de pacotes* (ou simplesmente *colisão*).

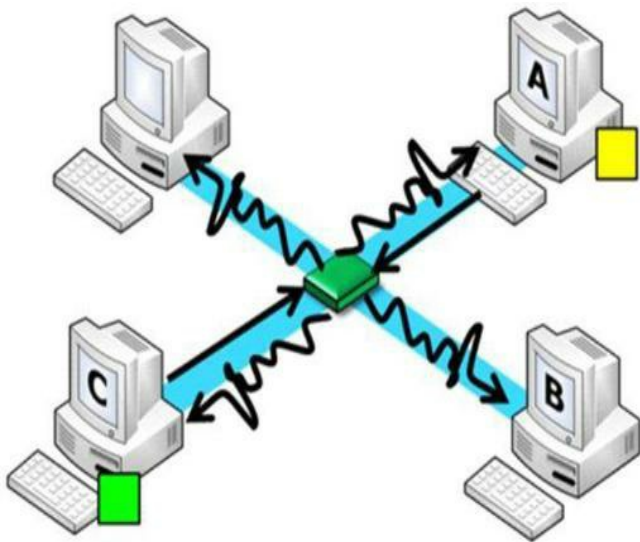


Figura 8.29 – Micros “A” e “C” tentaram transmitir ao mesmo tempo. “Deu Pau”! Colisão!

Lembre-se: em uma rede Ethernet, dois (ou mais) computadores podem tentar transmitir seus pacotes ao mesmo tempo (ou seja, têm condições de acessar o meio simultaneamente), mas isso é prejudicial – se isso acontecer, haverá colisão. A transmissão simplesmente não se completa quando há colisões.

Só há transmissão efetiva quando o pacote chega ao destino; portanto, em uma rede Ethernet, apenas um computador poderá transmitir pacotes por vez. (Isso só vai acontecer quando não houver colisões; em suma, quando aquele computador for, naquele dado momento, o único a ter

tido acesso ao meio.)

“Ei, João... Tá bom! Se a colisão é tão ruim assim, como é que a tecnologia Ethernet se livra dela? É possível ‘não haver colisões’ em uma rede Ethernet?”

É isso que está descrito das duas letras finais da sigla CSMA/CD.

- **CD (Collision Detection – Detecção de Colisões):** as colisões são inevitáveis? Sim! Para o CSMA/CD, que utiliza uma política reativa (corrigir) em vez de proativa (impedir). Outros protocolos (veremos adiante) conseguem prevenir colisões.

Mas se não dá para evitá-las, é necessário ter conhecimento quando uma delas acontece e, claro, realizar a correção necessária em tempo hábil. E é realmente engraçado como as placas de rede Ethernet corrigem as colisões ocorridas.

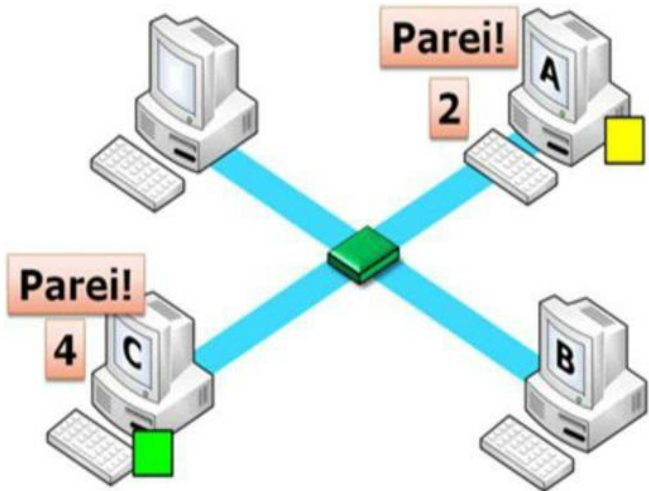
Vamos do início: quando dois computadores enviam sinais (pacotes) pela rede Ethernet, esses pacotes não chegam ao destino porque a colisão acontece. Ou seja, o micro “A” envia seus sinais e, ao escutar imediatamente algo que é bem diferente do que ele havia mandado (os sinais “misturados” com os de outro micro), “conclui” que a colisão ocorreu.

É fácil: se envie AAAAAAAAAA e escutei AxqH@@@#jKl é porque tem algo errado, não é mesmo?

E a outra estação (com a qual a colisão aconteceu) também sentirá a colisão e ambas (no caso da figura a seguir, “A” e “C”) simplesmente “param” – elas imediatamente param de transmitir (suspendem as transmissões dos bits subsequentes). Acontece o chamado **backoff**.

(É como se elas se “afastassem” com os braços levantados dizendo “ermão, na boa, parô, parô”.)

Ao pararem por causa da colisão, as duas (ou mais) estações envolvidas na “batida” resolvem as coisas por si só, sem chamar a polícia de trânsito ou o pessoal do seguro. Elas se sorteiam um número aleatório.



*Figura 8.30 – O início da solução do impasse causado pela colisão (backoff).*

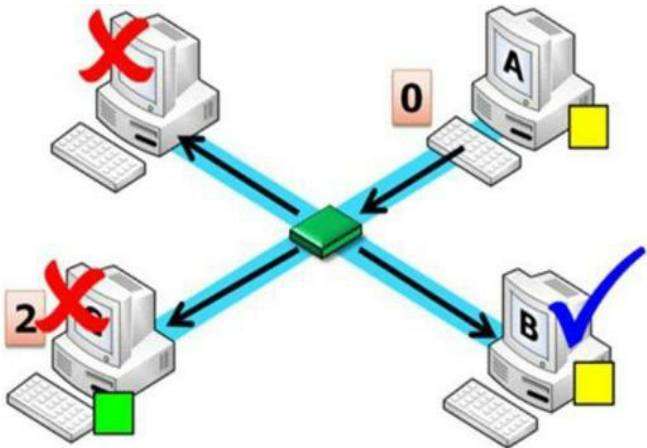
“João, para que esse número aí?”

Fácil: esse é o “atraso” (o tempo de backoff) – é o tempo que aquela estação vai aguardar para transmitir novamente seu pacote. É fácil notar, portanto, caro leitor, que a estação que sortear para si o menor número será a que vai transmitir seu pacote. À outra resta tentar novamente quando a “estação vencedora” terminar de transmitir seu pacote.

“Mas, João, se levarmos em consideração a figura anterior, a estação A vai transmitir daqui a 2, sei lá, milissegundos e a estação C vai transmitir daqui a 4 milissegundos – ou seja, daqui a 4 ms elas vão colidir novamente!”

Engano seu, nobre aluno. Quando passarem os 2 ms (não é exatamente essa a ordem de grandeza do “atraso”) de espera da estação “A”, ela vai escutar a rede (CS – Escutar a portadora) e vai ver que está livre. Portanto, o “meio físico” é todo dela. Quando chegarem os 4 ms, a estação “C” vai começar a tentar transmitir (fazendo o quê?) escutando a rede também. E, na pior das hipóteses, vai notar que a rede está sendo usada por alguém (a estação “A”).





*Figura 8.31 – Transmissão da estação “A” ocorrendo normalmente.*

Neste caso, a estação “C” vai se conformar e esperar que a rede esteja livre novamente. É igualzinho ao banheiro da casa em dia de feijoada. Quem conseguiu “transmitir” vai ficar lá até transmitir o “pacote todo” – os demais têm de esperar (até notarem que o banheiro – a portadora – está disponível).

“Ei, João, pensei num negócio engraçado: o número é aleatório, não é? E se as duas estações sortearam para si o mesmo valor para o atraso? Colidem de novo, não é?”

Sim! Se as duas estações sortearam para si o mesmo número (ou dois números muito próximos), elas vão esperar o mesmo tempo e transmitir simultaneamente (ou quase simultaneamente, que dá no mesmo), caracterizando outra colisão. Então, será necessário sortear novamente para corrigir a bendita “segunda colisão”.

Já vi uma cena hilária em trânsito que se assemelha muito com esse exemplo: dois carros bateram em um viaduto movimentado de Recife e as duas pessoas resolveram pacificamente a incômoda situação, mas tinham de tirar os veículos do meio do viaduto para liberar o tráfego. Pois imaginem só: durante a “manobra” para solucionar a batida e livrar o tráfego, eles bateram novamente. Vá entender...

Se houver uma segunda colisão, as estações são obrigadas a sortear um número de backoff em um intervalo duas vezes maior. Como o intervalo de escolha está maior, a probabilidade de que as estações sorteiem números iguais (ou muito próximos) diminui. Vamos ver isso de uma forma

mais fácil de entender:

1. Primeira colisão: as estações sorteiam tempos entre 0 e 2ms; (escolheram 1,32 e 1,32).
2. Segunda colisão: as estações sorteiam atrasos entre 0 e 4ms; (por exemplo, escolhem 3,2 e 3,2).
3. Terceira colisão: as estações sorteiam atrasos entre 0 e 8ms; e assim por diante.

Esse aumento do atraso é chamado de Backoff Exponencial Binário (binário porque o aumento é feito multiplicando-se o intervalo anterior por 2) e tem a função de diminuir, a cada mudança, a probabilidade de ocorrência de novas colisões.

Bem, é isso... Essas são as regras das redes Ethernet. Simples, não? Vamos, então, resumi-las em algumas linhas:

1. Uma estação que deseja transmitir algo fica “escutando a rede” para determinar quando poderá transmitir. Enquanto essa rede estiver sendo usada, aquela estação não se sente apta a transmitir. É fácil detectar se a rede está sendo usada (ou seja, se há outra transmissão atualmente acontecendo): é só a placa de rede ficar “de olho” no que acontece nos pinos 3 e 6 de seu conector – os pinos de recepção de dados, como já havíamos visto.
2. Se a estação detecta a ausência de sinal nos pinos 3 e 6, quer dizer que a rede “calou-se”, ou seja, esta é a deixa: a estação envia seus sinais (pacote) na forma de pulsos elétricos pelos pinos 1 e 2 de seu conector.
3. Depois de enviar seus sinais para a rede, duas coisas podem acontecer:
  - 3a. se aquela estação que transmitiu era a única a utilizar o meio físico naquele momento, então é fácil concluir que os sinais elétricos serão devidamente transmitidos em broadcast (para todas as estações) e que quando atingir a estação de destino adequada, esses sinais serão devidamente recebidos, concretizando o envio bem-sucedido do pacote (claro que é bom lembrar de que as demais estações vão rejeitar aquele pacote porque não é direcionado a elas); ou
  - 3b. se aquela estação não é a única a ter acesso ao meio (tentar transmitir), seus sinais vão se misturar com os da(s) outra(s) estação(ões) transmissora(s), formando um conjunto de dados estranhos e deturpados, caracterizando uma colisão.
4. Se houver colisão, as estações envolvidas param de transmitir seus sinais e imediatamente procedem com um “sorteio” de um número aleatório. Esse número determina quanto tempo aquela estação vai esperar para transmitir novamente (atraso, ou backoff). O intervalo de escolha do valor do atraso é aumentado exponencialmente a cada nova colisão.

#### **8.6.1.3. Conclusões sobre a arquitetura Ethernet**

Nas redes Ethernet, as colisões acontecem! Isso é um fato! Outro fato é que a rede Ethernet “tem consciência” de que elas acontecem, e prevê correções para elas. A conclusão a que chegamos com isso não é tão positiva: se as placas de rede vão “gastar” certo tempo corrigindo colisões, a rede Ethernet não alcança as velocidades descritas efetivamente.

Mesmo em uma rede Fast Ethernet (que diz ter velocidades de 100 Mbps), a quantidade de colisões grande pode diminuir em muito a efetiva velocidade da rede. Ou seja, em vez de transmitir dados a 100 Mbps, a rede vai transmitir, na realidade, a 40, 50, 60 Mbps – dependendo da quantidade de colisões.

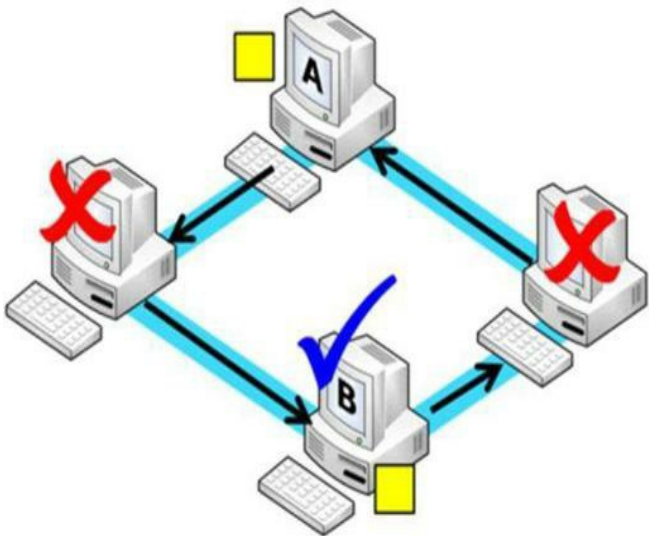
E a quantidade de colisões está diretamente ligada à quantidade de tráfego, que está diretamente associada à quantidade de micros na rede. Então, cheguei aonde queria: “Quanto mais estações (micros) em uma rede barramento (como é o caso da Ethernet), mais lenta a rede será” – porque vai passar grande parte do tempo corrigindo as besteiras que acontecem na rede em vez de transmitir dados novos.

E assim chegamos ao final da explicação sobre as redes Ethernet, com detalhamento na sua “cartilha de boas maneiras” – o protocolo CSMA/CD. Vamos a uma tecnologia de redes já antiga (e não mais utilizada).

#### **8.6.2. Token Ring (IEEE 802.5)**

A Arquitetura Token Ring foi desenvolvida pela IBM para ligar computadores em anel e hoje é regulamentada pela norma 802.5 do IEEE. A Arquitetura Token Ring já foi muito mais utilizada, mas hoje perdeu completamente seu espaço para as redes Ethernet. A taxa de transferência máxima de uma rede Token Ring é de 16 Mbps (um pouco mais que o Ethernet original, mas com certeza bem menos que o Fast e o Gigabit Ethernet).

Na arquitetura Token Ring, as placas de rede dos computadores têm comportamento ativo, ou seja, elas funcionam como o que chamamos de repetidores. Para que uma mensagem atravessasse todo o anel, ela deverá passar por todas as estações, que, por sua vez, irão receber os sinais elétricos e retransmiti-los para os demais computadores (na verdade, é a placa de rede Token Ring que faz isso). É preciso lembrar também que a mensagem chega ao destino e retorna para a origem. A mensagem atravessa todo o anel.



*Figura 8.32 – Uma transmissão na rede Token Ring de “A” para “B”.*

É bastante simples o funcionamento da rede Token Ring:

1. Um micro envia dos dados pelo anel.
2. A mensagem (pacote) atravessa todos os computadores do anel, sendo passada adiante por estes se não forem o destinatário da mensagem.
3. O micro destino recebe o pacote (copia-o para si) e o passa adiante.
4. A mensagem retorna ao computador que a enviou, com isso, este poderá transmitir seu próximo pacote ou liberar a rede para ser usada para a transmissão de outro pacote, vindo de outra estação.

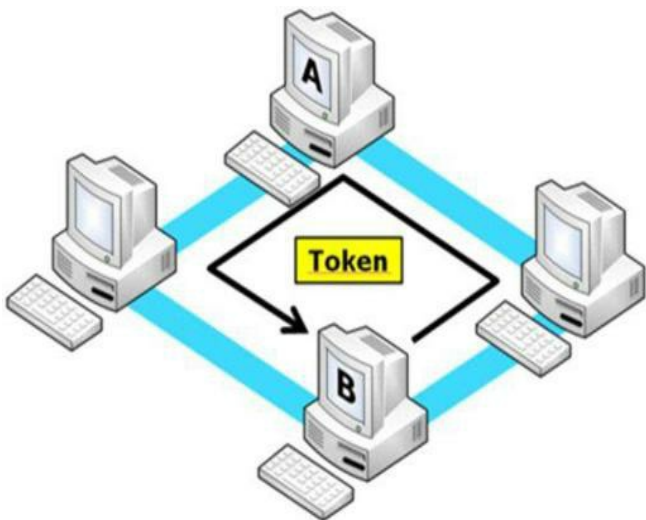
Mas, até aqui, nenhuma novidade, não é mesmo? Pois esse é o funcionamento de uma rede anel, como já havíamos visto. A questão principal é: como o micro “A” obteve o direito de falar na rede? Ou, em palavras mais técnicas: como uma estação obtém o acesso ao meio físico?

#### **8.6.2.1. Quadro Token (permissão)**

Como um computador obtém acesso para transmitir dados em uma rede Token Ring? Simples: há um mecanismo que controla o fluxo permitindo que apenas um computador obtenha acesso ao meio por vez. Esse mecanismo é chamado quadro *token* (também chamado de *permissão* ou *ficha*).

Na verdade, um token é apenas um pacote que *não transmite dados* e que fica circulando o anel o tempo todo. Portanto, o método de acesso da rede Token Ring é chamado *Passagem de Permissão*, *Passagem de Token* ou *Passagem de Ficha*.

Então, vejamos: o token (um pequeno quadro sem dados significativos) fica trafegando pela rede de micro em micro (a rede é anel, lembre-se!). Enquanto o token estiver passeando pela rede, nenhuma estação está transmitindo nada.



*Figura 8.33 – O token trafegando pela rede (isso é sinal de que a rede está livre).*

Quando um computador deseja transmitir seus dados pela rede, este “captura” o token e, de posse dele, envia um pacote (sim, apenas um). Quando essa estação obtiver o pacote de novo



microfone é o token. É ele que dá permissão (e condição) aos participantes de transmitirem suas opiniões.

Uma característica que vale a pena ressaltar quanto ao modelo de passagem de permissão é que, com ele, não há colisões de pacotes. Verdade! Não há! Porque não há acesso múltiplo, pois apenas um computador (aquele que detém o token) vai conseguir acesso ao meio físico. Isso gera homogeneidade na velocidade da rede, mesmo com quantidades variadas de estações ligadas a ela.

Outra coisa muito interessante com relação às redes Token Ring é o fato de que, mesmo funcionando em anel, essa rede é fisicamente conectada em estrela (o que, aliás, é característica de todas as redes anel), utilizando-se de equipamentos concentradores conhecidos como MAU (Multistation Access Unit) – que popularmente são chamados de “hubs Token Ring”.



*Figura 8.35 – Uma MAU (um concentrador Token Ring).*

Uma MAU é um dispositivo capaz de conectar fisicamente os computadores, criando internamente o anel entre elas. Uma MAU também consegue detectar falhas em uma estação e, de imediato, separá-la do restante do anel (fechar o anel sem ela) para que a rede não pare de funcionar.

Bom, mas chega de velharia! Vamos analisar agora uma das mais novas e interessantes arquiteturas de redes da atualidade: a Wi-Fi.

### **8.6.3. Wi-Fi (IEEE 802.11) – Redes LAN sem fio**

Como o nome já diz, esta arquitetura de rede não utiliza cabos de cobre nem fibra óptica. Os sinais são transmitidos entre os computadores através de ondas eletromagnéticas.

Wi-Fi é, portanto, uma arquitetura que especifica o funcionamento de uma WLAN (Wireless LAN, ou LAN sem fio). Note que WLAN é um termo genérico, pois significa qualquer “rede local sem fio”, porém Wi-Fi é o termo que designa essa tecnologia, também conhecida como 802.11. (Porque essa arquitetura de redes foi padronizada segundo a norma 802.11 do IEEE.)

Na verdade, Wi-Fi significa Wireless Fidelity (ou Fidelidade sem fio) e é um “título” dado a todos o equipamentos (e programas) que “seguem à risca” a cartilha proposta pelo padrão IEEE

802.11. Portanto, se um equipamento mereceu o título de Wi-Fi, é sinal de que ele é perfeitamente compatível (ou seja, está em concordância) com os padrões descritos para redes locais sem fio.

As redes no padrão 802.11 usam uma topologia lógica de barramento (portanto, trabalham por difusão) e controlam o acesso dos computadores através de um sistema semelhante ao CSMA/CD das redes Ethernet. Nas redes 802.11, o método de acesso ao meio é chamado **CSMA/CA** (Carrier Sense with Multiple Access and **Collision Avoidance** – algo como Sensor de Portadora com Acesso Múltiplo **Evitando Colisões**).

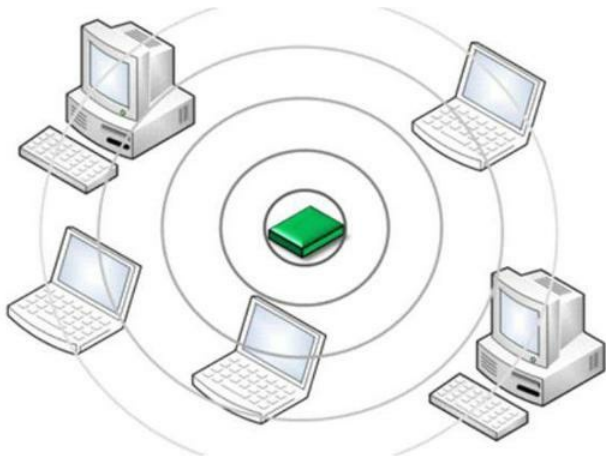


Figura 8.36 – Funcionamento da Rede IEEE 802.11 em modo Infraestrutura.

Nessa rede, os computadores são dotados de placas de rede especiais, criadas apenas para essa finalidade. São placas de rede que possuem antenas para transmitir e receber os sinais das outras placas em vez de conectores como o RJ-45.

Uma rede Wi-Fi pode ser montada basicamente de duas maneiras:

- **Modo Infraestrutura:** os micros são ligados entre si por meio de um equipamento central (algumas vezes chamado de hub sem fio). Esse equipamento recebe as transmissões de uma estação e as passa para todos (difusão). Esse equipamento é chamado de **Ponto de Acesso** (**Access Point**);



- **Modo Ad-Hoc:** os micros são ligados diretamente uns aos outros (placa de rede direto para placa de rede), ou seja, sem a presença de um ponto de acesso.

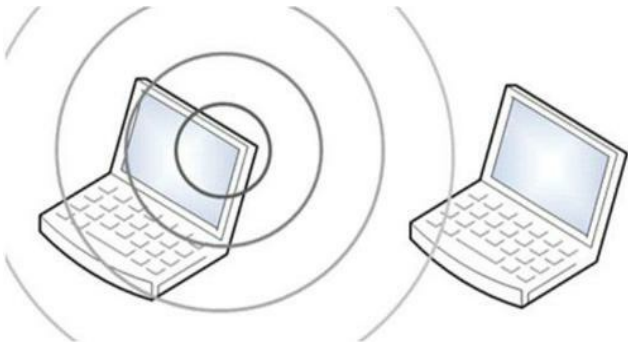


Figura 8.37 – Rede Wi-Fi em modo Ad-Hoc.

#### 8.6.3.1. Subpadrões 802.11

Dentro do padrão IEEE 802.11, há diversos subpadrões desenvolvidos e incentivados por várias empresas, entre eles podemos destacar quatro que são diferentes na frequência que utilizam para transferir os dados e na taxa máxima de transferência.

- **802.11b:** o padrão mais antigo. Os equipamentos que trabalham neste padrão usam uma frequência de **2,4 GHz** e transmitem dados a **11 Mbps** (pouco mais que a velocidade da arquitetura Ethernet original).
- **802.11g:** atualmente, é o padrão de rede Wi-Fi mais usado. Também utiliza a faixa de frequência dos **2,4 GHz** (o que garante a perfeita comunicação entre equipamentos “b” e “g”). Transmite dados a **54 Mbps**. É claro que para transmitir a 54 Mbps, é necessário que todos os equipamentos envolvidos sejam do padrão “g”.
- **802.11a:** é um padrão pouco usado no Brasil que utiliza a faixa de frequência de **5 GHz** para transmitir a **54 Mbps**. Devido à diferença de frequência, equipamentos nesse padrão não conseguem se comunicar com os outros padrões citados.
- **802.11n:** realiza transmissões da ordem de **300 Mbps** (três vezes mais que o Fast Ethernet), usando as duas faixas de frequência possíveis (**2,4 GHz e 5 GHz**) para que os equipamentos “n” possam se comunicar com outros de todos os padrões.

Alguns fabricantes criaram equipamentos “n” com velocidades de até 600 Mbps, mas que só

funcionam se todos os equipamentos envolvidos (placas de rede e pontos de acesso) forem da mesma marca.

- **802.11ac:** mais recente dos padrões em comercialização (para se ter uma ideia, o primeiro equipamento “ac” do mercado foi lançado em setembro/2012). Esta variante do 802.11 admite velocidades de mais de **1 Gbps (1.000 Mbps)**, usando frequência de **5G Hz**.

Este padrão ainda é **draft** (rascunho), ou seja, ainda não foi totalmente aprovado pela equipe 802.11. Estima-se para o final de 2013 a aprovação deste padrão. Mesmo sendo draft (“projeto de lei”, se for mais fácil entender assim), já há fabricantes construindo equipamentos seguindo as determinações deste padrão.

- **802.11ad:** padrão em desenvolvimento. Admitirá velocidades de até 7 Gbps. Não se pode dizer nada concreto acerca deste padrão, porque, daqui para sua aprovação (estima-se em 2015), muita coisa pode mudar!

Na hora de comprar equipamentos para montar sua rede Wi-Fi, é bom verificar com cuidado o padrão que deseja usar (e verificar se todos os equipamentos adquiridos estão em concordância com aquele padrão), para que se obtenha o melhor resultado de desempenho.

Lembre-se, também, de que essa velocidade de transmissão é atingida de acordo com a distância entre as estações e o estado do “ambiente” (se está livre, se apresenta obstáculos etc.). Quanto mais longe uma estação estiver de outra ou de um ponto de acesso, mais lenta será a transmissão para essa estação.

Ou seja, se uma estação estiver muito próxima a um ponto de acesso, em uma rede “g”, eles vão se comunicar a 54 Mbps realmente. Mas se houver certa distância entre eles e/ou obstáculos diversos, como paredes ou reservatórios de água (como aquários ou garrações de água mineral – que, diga-se de passagem, são bem prejudiciais), essa velocidade vai cair para 48, 36, 24, 22, 18, 11, 5, 2 e até 1 Mbps. Essa diminuição de velocidade gradativa em função da distância das estações e dos empecilhos entre elas é chamada de fall-back.

### 8.6.3.2. CSMA/CA

De forma análoga ao CSMA/CD, neste método de acesso, os computadores têm o direito de acessar o meio (no caso, o “ar”), apenas analisando a portadora (o “direito” de eles transmitirem seus sinais).

A diferença está no fato de que os computadores, antes de enviar seus dados, mandam um sinal de “vou enviar em tantos milissegundos, ok?”, e isso alerta aos demais computadores para não enviarem concomitantemente. Portanto, como o já foi dito, o “CA” da sigla significa “Collision Avoidance”, ou “Evitar Colisões”.

### 8.6.4. Segurança nas redes Wi-Fi

Se tem um “calcanhar de Aquiles” nas redes Wi-Fi, este é, sem dúvidas, a parte da segurança. Esse critério é sempre visto de forma “atravessada” quando se fala em redes 802.11.

Quando uma rede Wi-Fi está “aberta” (sem proteção de qualquer tipo), entrar nela é um processo relativamente fácil – basta que a estação “invasora” esteja dentro do raio de cobertura da antena do ponto de acesso e irá captar os sinais daquela rede e, em alguns instantes, fará parte dela como uma de suas estações legítimas.

Como alguns pontos de acesso têm potência de sinal muito alta, um invasor poderia, digamos, entrar na rede de uma empresa usando um laptop dentro do carro estacionado próximo ao prédio. (E uma antena especial mais forte, às vezes feita, pasme leitor, de uma lata de batatas fritas – daquelas latas cilíndricas revestidas de papel alumínio internamente.)

“Lata de batatas fritas, João? daquelas que compramos em qualquer supermercado?”

Sim, essa mesmo! Não quis citar o nome porque pareceria merchandising! E ainda tem mais! A de cebola e salsa é a melhor para fazer a antena!

“Ai já é demais. Por que justo a de cebola e salsa? Existe algum componente especial na cebola ou na salsa?”

Não. Simplesmente porque, para usar a lata como antena, ela tem de estar vazia (e, adivinhe, a gente faz isso comendo as batatas!), e *eu adoro a de cebola e salsa*.



*Figura 8.38 – Antena direcional feita com lata de batatas.*

Só para explicar, o formato cilíndrico longo da lata e o revestimento metálico interno fazem dela uma perfeita antena direcional (antenas que captam/transmitem sinais basicamente na direção para a qual são apontadas).

Neste ponto (a segurança), as redes cabeadas (com fios) são muito melhores. Porque para ter acesso à rede com fio, a estação invasora tem de estar ligada fisicamente a um cabo da rede, o que, para quem está fora do prédio da instituição a ser invadida, é bastante difícil.

Para tornar as redes sem fio mais seguras, foram criados alguns métodos (na forma de recursos de software, como protocolos) para criptografar (embaralhar) os dados que trafegam pela rede para, em teoria, impedir o acesso das estações bisbilhoteiras.

#### 8.6.4.1. WEP (Wired Equivalent Privacy)

O protocolo WEP (que significa “Privacidade Equivalente à Cabeada”) foi o primeiro protocolo criado com a finalidade de permitir a criptografia dos dados dos pacotes antes de eles serem enviados pela estrutura da rede. Teoricamente isso embaralha os dados de tal forma que somente os computadores que conhecem o segredo (a chave da rede) tenham condições de se comunicar naquela rede.

O WEP usa o algoritmo RC4 para o processo de criptografia, podendo utilizar chaves de 40 bits ou 104 bits (a escolha é de quem configura a rede).

A principal vulnerabilidade no *modus operandi* do WEP é o fato de a chave ser sempre a mesma na comunicação (mudar a chave é um processo manual e exige que o administrador da rede faça-a em todos os micros) – portanto, a captura de alguns pacotes da rede (existem programas que fazem isso facilmente) e a análise desses pacotes permitiriam a descoberta da chave da rede mais cedo ou mais tarde (hoje em dia, é sempre “mais cedo”).

O fato é que já foram desenvolvidos softwares utilitários que capturam os pacotes e “descobrem” a chave da rede em questão de minutos (na verdade, em alguns casos, em menos de um minuto), portanto, o WEP é, declaradamente, *considerado inseguro*.

#### 8.6.4.2. WPA (Wi-Fi Protected Access)

Eis aqui o sucessor (com louvor) do WEP. No WPA, a criptografia é mais forte que no WEP e sua arquitetura de compartilhamento de chaves é mais consistente, reduzindo muito a possibilidade de quebra do segredo e consequente invasão da rede.

Dentre as melhorias do WPA em relação ao WEP estão:

- **Algoritmo de Criptografia TKIP:** o algoritmo TKIP (Temporary Key Integrity Protocol – ou Protocolo de Integridade de Chave Temporária) usa uma chave-base (chave inicial) de 128 bits que vai gerar chaves criptográficas mutáveis, associadas unicamente a cada pacote. Sim! Cada pacote tem uma chave diferente! Isso faz com que não se possa estudar um padrão mesmo que se capturam vários pacotes para analisá-los, tornando muito difícil a descoberta da chave-base.
- **Autenticação de Usuários (EAP – 802.1x):** em uma rede protegida por WPA, é possível restringir ainda mais o acesso aos recursos da rede por meio de um sistema que autentica usuários (por meio de nome, senha, certificado digital, biometria ou outro método). Com isso, antes de ter acesso à rede, a estação tem de provar que o usuário que a está utilizando é merecedor de estar conectado. O protocolo que faz esse tipo de verificação chama-se EAP (Extensible Authentication Protocol – Protocolo de Autenticação Extensível) e está descrito (protocolado, definido) no padrão 802.1x, que versa sobre segurança nas redes de modo

geral.

“Ei, João. Posso usar esses dois recursos em minha rede Wi-Fi?”

Na verdade, a Wi-Fi Alliance (a “instituição” que pesquisa, desenvolve e padroniza as tecnologias e processos para redes 802.11) diz que há dois tipos de WPA:

O WPA-Personal, para ser usado em redes caseiras ou de pequenas empresas, que funciona por meio de uma PSK (Pre-Shared Key – chave pré-compartilhada) – que é, nada mais nada menos, a nossa “chave-base” – fornecida manualmente pelo administrador da rede.

Ou seja, no modo WPA-Personal, alguém (de carne e osso) vai a cada micro legítimo da rede configurar manualmente a chave-base, que é chamada de passphrase (“Frase passe” ou “senha frase”) e consiste em um trecho alfanumérico (letras e números) que pode ter até 63 caracteres.

No modo WPA-Enterprise, não há PSK (não é necessário escrever a passphrase em todos os micros), mas, em compensação, os micros devem usar um servidor de autenticação para enviar suas credenciais (nome, senha, certificado digital etc.) para que a chave inicial seja criada e enviada de volta às estações. Portanto, o uso da segurança do 802.1x é restrito ao modo WPA-Enterprise.

“João, a diferença entre o Personal e o Enterprise é, pelo que pude perceber, somente na forma como essa tal de chave inicial é criada, não é mesmo?”

Precisamente! Depois de criada a chave-base (seja pré-fornecida, ou pela autenticação em um servidor apropriado), o WPA seguirá seu caminho com o TKIP, criando chaves de criptografia para cada pacote a ser transmitido.

#### 8.6.4.3. WPA 2 (IEEE 802.11i)

O protocolo WPA já atingiu uma nova versão: o WPA2, que é totalmente compatível com as especificações contidas na padronização IEEE 802.11i. (Uma norma técnica que padroniza a segurança nas redes Wi-Fi.) O WPA2 apresenta, como protocolo de criptografia, o AES (padrão internacional de criptografia simétrica) com chaves maiores (de 256 bits) e oferece todos os recursos já existentes na primeira versão do WPA.

Ao adquirir qualquer equipamento Wi-Fi (pontos de acesso, placas de rede sem fio, roteadores sem fio), é bom verificar as compatibilidades com relação à segurança.

#### 8.6.5. Mais glossário Wi-Fi

##### 8.6.5.1. MIMO (Multiple-Input, Multiple-Output)

MIMO (Múltiplas Entradas, Múltiplas Saídas) é uma tecnologia que aumenta consideravelmente a velocidade e o alcance das redes Wi-Fi. Usado inicialmente nas redes “g” (como opcional), essa tecnologia está amplamente difundida nos equipamentos do padrão “n” (802.11n), pois as redes “n” se baseiam nela.

Consiste em utilizar o fenômeno da propagação múltipla das ondas eletromagnéticas em diferentes ângulos para aumentar a capacidade de transmissão e recepção de dados simultaneamente através do uso de *múltiplas antenas* tanto nas placas de rede quanto nos pontos de acesso.



*Figura 8.39 – Roteador + Ponto de acesso com tecnologia MIMO.*

As redes “n”, “ac” e “ad” e todas as versões futuras do 802.11 têm no MIMO uma importante característica para aumentar suas velocidades originais de transmissão.

#### **8.6.5.2. Hotspot**

Hotspot designa um local público onde é possível (por meio de pagamento ou não) acessar a Internet através de uma rede Wi-Fi. Há hotspots em hotéis, restaurantes, aeroportos, mercados municipais (em Porto Alegre, por exemplo), faculdades e até mesmo hospitais.

#### **8.6.5.3. SSID**

É, tão somente, o “nome” da rede Wi-Fi. Quando um usuário inicializa sua placa de rede Wi-Fi, seja um laptop ou um micro de mesa, a placa começa a captar todos os sinais das redes próximas a ela.

Para saber a qual rede se conectar, é necessário saber o SSID (Service Set Identifier – Identificador de Conjunto de Serviços) dessa rede. É o conjunto de caracteres (letras e números) que identifica uma rede e a diferencia das demais.

Veja o programa gerenciador de conexões sem fio do Windows Vista mostrando as redes disponíveis ao alcance dele.



Figura 8.40 – Cinco redes Wi-Fi detectadas (conectado à Família).

Bem, leitor: com isso terminamos a análise dos principais tipos de redes locais de computadores (LAN) e agora partiremos para o estudo das redes de maior alcance de área.

### 8.7. Arquiteturas para Mans e Wans

Há também tecnologias (arquiteturas) importantes de ser estudadas em redes metropolitanas (MAN) e redes de longo alcance (WAN). Essas arquiteturas permitem a comunicação entre computadores distantes entre si alguns (ou muitos) quilômetros.

São utilizadas por várias empresas, desde pequenas que queiram ligar suas filiais em uma mesma cidade até gigantes de telecomunicações que queiram expandir seu “domínio” e oferecer uma estrutura mais ampla para seus assinantes.

Eis algumas tecnologias importantes que devemos conhecer. (Não que apareçam assim nas provas o tempo todo, mas é possível que sejam citadas em uma ou outra)

#### 8.7.1. ATM

ATM (Asynchronous Transfer Mode – Modo de Transferência Assíncrono) é uma tecnologia de comunicação de dados que permite a construção de redes LAN, MAN e WAN. O ATM é uma arquitetura de rede orientada a conexão, ou seja, antes de mandar o primeiro pacote de dados, o emissor verifica se a conexão entre ele e o receptor foi estabelecida (essa conexão é chamada “circuito virtual”).

A principal proposta desta arquitetura é permitir o tráfego de vários tipos de dados: voz, vídeo, serviços de rede etc. Pode-se atingir 155 Mbps (em cabos de cobre ou fibra óptica) ou até 622

Mbps (usando exclusivamente a fibra óptica).

Uma das principais características da rede ATM é a forma como ela transfere os dados. Diferentemente de várias outras redes, que usam blocos de dados enormes (e com tamanhos variados), a rede ATM divide os dados a serem transmitidos em pacotes muito pequenos (conhecidos como *células*). Uma célula ATM tem exatamente **53 Bytes**, dos quais **5 são para cabeçalho** (informações de endereçamento e caminho para a entrega dos dados) e **48 são de dados** propriamente ditos (payload).

Atenção: esta é a principal característica que se pode cobrar sobre o ATM: o tamanho de sua célula (pelo menos, historicamente, é o que foi cobrado!).

Por causa do nome de seus pacotes (células), o ATM é conhecido como **Cell Relay** (algo como “chaveamento de células”).

Essa tecnologia está sendo amplamente usada nas operadoras de telecomunicações, como as empresas telefônicas, para a interligação entre suas centrais regionais e até mesmo em alguns serviços de ADSL (Internet Banda Larga) para usuários finais.

### 8.7.2. Frame Relay

Frame Relay é uma tecnologia para ligação de computadores em WAN descendente da antiga tecnologia X.25. No Frame Relay, os dados são separados em unidades conhecidas como frames (quadros) que são enviados através de linhas que transmitem sinais analógicos.

Essa tecnologia é usada (ainda) por empresas de telecomunicações (como as operadoras telefônicas) para permitir a ligação com centrais e usuários longe dos centros, onde tecnologias como ATM ou ADSL não podem chegar – como em áreas rurais, por exemplo.

As operadoras que fornecem o serviço de Frame Relay o vendem em várias velocidades, desde 56 Kbps a 1,5 Mbps (para usuários finais) até as taxas de transmissão mais altas, usadas para grandes clientes e interligação entre centrais da própria operadora (até 100 Mbps). Mas essa tecnologia está caindo em desuso graças ao ATM e a outras tecnologias novas para WAN.

### 8.7.3. WiMAX (IEEE 802.16)

WiMAX (Worldwide Interoperability for Microwave Access ou Interoperabilidade Mundial para Acesso por Micro-ondas) é uma tecnologia de transmissão de dados para redes de computadores de área metropolitana (MAN) sem fio. Daí o nome de WMAN (Wireless MAN – MAN sem fio).

O padrão 802.16 foi totalmente homologado em 2002 e hoje já é realidade em algumas cidades do mundo (incluindo algumas aqui no Brasil, a exemplo de Belo Horizonte e Rio de Janeiro).

Através do WiMAX, uma antena é colocada em um determinado ponto da cidade e esta cria uma área de cerca de 50 km de raio. A velocidade praticada por essa tecnologia chega a 70 Mbps (pouco mais que as redes Wi-Fi “a” e “g” e menos que a “n”).

O WiMAX usa uma faixa de frequência de 2,3 a 2,5 GHz e, em alguns países, de 3,3 GHz. Futuras aplicações dessa tecnologia darão conta de uso de outras faixas de frequência (algumas superiores a 10 GHz).

Não se esqueça disto: WiMAX é uma tecnologia para redes MAN (metropolitanas) e um uso



interessante para essa tecnologia é o fornecimento de Internet em banda larga para locais onde as operadoras telefônicas e de TV a cabo não podem ir para fornecer alta velocidade no acesso à Internet.

#### 8.7.4. IEEE 802 – redes de computadores

O IEEE (Instituto de Engenheiros Elétricos e Eletrônicos) é um órgão que “dita as regras” acerca de quase todos os equipamentos de informática e telecomunicações atualmente, como já sabemos. Algumas poucas são exceções ao “juízo” do IEEE.

Como vimos exaustivamente, o IEEE também escreveu padrões (documentos de padronização) para quase todos os tipos de tecnologias de redes de computadores atualmente vigentes. Esse conjunto de normas (feito especialmente para determinar os padrões relacionados com as redes de computadores) é chamado de *Projeto 802*.

Não é nenhuma “superstição”, não. O grupo de trabalho que padronizaria as redes de computadores foi fundado em fevereiro (mês 2) de 1980. (Isso não vai cair na prova... É só para você saber... Cultura inútil.)

O Projeto IEEE 802 é dividido em diversos WG (Working Groups – Grupos de Trabalho) e cada um desses grupos atua em um cenário diferente dentro do vasto mundo das redes de computadores.

São mais ou menos como as “comissões parlamentares” na Câmara e no Senado, que são grupos de “trabalho” que dividem-se para tratar de assuntos específicos, como a (faz-me rir) Comissão de Ética, como a Comissão de Segurança, a Comissão de Educação etc.

Nota: em respeito aos Grupos de Trabalho do Projeto 802, gostaria de deixar claro que este último parágrafo é apenas a título de analogia, quanto à divisão dos grupos em projetos menores. Nenhuma comparação quanto à utilidade, à ética ou à efetividade dos Grupos foi feita com relação às comissões parlamentares, até mesmo porque eu sei que os Grupos funcionam!

Não sei bem o porquê (pois acho que não são cobrados), mas resolvi listá-los (pelo menos alguns deles) aqui:

| O Grupo... | ... Trata de...                                    |
|------------|----------------------------------------------------|
| 802.1      | Padrões gerais para LAN e MAN, incluindo segurança |

|        |                                                      |
|--------|------------------------------------------------------|
|        | em Redes de Computadores                             |
| 802.2  | Logical Link Control (parte da Camada 2) (já era!!)  |
| 802.3  | Redes LAN Ethernet                                   |
| 802.4  | Redes LAN Token Bus (já era!!!)                      |
| 802.5  | Redes LAN Token Ring (já era!!!)                     |
| 802.11 | Redes LAN sem fio (WLAN – Wireless LAN) – Ex.: Wi-Fi |
|        | Redes PAN sem fio (WPAN – Wireless                   |

802.15

PAN) – Ex.:  
Bluetooth

802.16

Redes MAN sem  
fio (WMAN –  
Wireless MAN) –  
Ex.: WiMAX

802.20

Redes Móveis de  
Banda Larga sem  
fio (MBWA) –  
redes para  
comunicação móvel  
(em veículos, por  
exemplo) em banda  
largaInteroperabilidade  
entre diversas  
tecnologias de  
redes sem fio

802.21

(802.11, 802.16,  
802.20 e  
tecnologias não  
padronizadas em  
802)

Então é isso. Além de “meterem o nariz” para padronizar tudo, a gente ainda tem de decorar essas doidices deles (pelo menos, talvez um dia caia em prova).

Bem, leitor, com isso terminamos as tecnologias de redes para MAN e WAN. Vamos agora analisar os equipamentos que podem ser encontrados em redes diversas.

## **8.8. Equipamentos usados nas redes**

Em muitas provas é comum exigir o conhecimento nas características principais de alguns equipamentos usados em redes, como os que são mostrados a seguir.

### **8.8.1. Placa de rede (ou adaptador de rede)**

É o equipamento que deve existir em cada computador para que eles possam se conectar a uma rede local (LAN). A placa de rede (ou NIC – Network Interface Card, – Placa de Interface de Rede, ou ainda Adaptador de Rede) é um periférico normalmente instalado no interior do gabinete do computador, diretamente em um dos slots da placa-mãe (normalmente um slot PCI).

Também é possível que a placa de rede já seja fabricada na própria placa-mãe (prática, aliás, muito comum hoje em dia) tanto nos notebooks quanto nos micros de mesa (desktops) – é a chamada placa de rede on-board, como já foi visto no capítulo sobre hardware.

Uma placa de rede é fabricada para se comunicar com um tipo específico de arquitetura, ou seja, com um determinado tipo de protocolo, cabeamento também específico entre outras coisas. Logo, há vários tipos de placas de rede disponíveis no mercado, pois há vários tipos de arquiteturas de redes. (Lembre-se de que as duas mais usadas são a Ethernet e a Wi-Fi.)

“Ô, João, quer dizer que um computador pode ter mais de uma placa de rede? Uma para cada arquitetura de redes na qual ficará ligado?”

Precisamente, mas não só isso! Um computador pode ter mais de uma placa de rede de mesma arquitetura (nada impede, a não ser a falta de sentido disso na maioria dos casos). Um exemplo bem simples são os laptops (ou notebook, se preferir) vendidos atualmente: todos eles saem das fábricas com duas placas on-board – uma placa Ethernet e outra placa Wi-Fi.

Veja dois exemplos de placas de rede conectáveis ao barramento PCI das placas-mãe dos micros desktop (micros de mesa).



*Figura 8.41* – Placa de Rede Ethernet (note o conector RJ-45).



*Figura 8.42 – Placa de Rede Wi-Fi (note a antena) “G” encaixável no barramento PCI.*

“Ei João, mas as duas placas mostradas não servem para micros portáteis, não é mesmo? Afinal, como elas vão ficar encaixadas?”

Perfeito! Qualquer placa de expansão (lembra do nome, né?) que se conectar aos barramentos PCI, AGP, PCI Express X1 ou X16 não é conectável em um laptop, pois não há slots dessas interfaces neles.

Em um micro portátil, praticamente todas as placas são instaladas na própria placa-mãe, ou seja, são todas on-board. Em alguns casos, pode-se comprar placas especiais de expansão que encaixam na interface PCMCIA (CARD BUS) – que hoje é menos comum – ou pequenos adaptadores que são plugados em qualquer porta USB, como o visto a seguir.



Figura 8.43 – Adaptador Wi-Fi USB – pode ser usado em laptops e desktops.

#### 8.8.1.1. Endereço MAC (endereço físico)

Cada placa de rede que é fabricada recebe um número único, que a diferencia de qualquer outra placa. Esse número é conhecido como *MAC Address* (Endereço MAC) ou *Endereço Físico*.

O endereço MAC é uma espécie de “número de chassi” da placa de rede, pois cada fabricante coloca o endereço no momento da montagem da placa e esse endereço não será usado por nenhuma outra placa de rede no mundo.

O endereço MAC é formado por 48 bits (48 “zeros e uns”). Isso significa que o endereço MAC é, na verdade:

**100000010000001001000101001110110010101000001110**

Mas normalmente, o endereço MAC de uma placa de rede é representado (e visto por nós, humanos) como um conjunto de seis duplas de dígitos hexadecimais. A conversão de números binários para hexadecimais e vice-versa é um procedimento relativamente simples e é abordado na última parte deste livro. Eis o mesmo endereço MAC, desta vez em hexadecimal:

**81:02:45:3B:2A:0E**

Como os endereços MAC são gravados nas memórias ROM das placas de rede, eles não

podem ser alterados e estão, para sempre, associados àquela placa de rede em si (àquele exato equipamento).

“Ei, João, vai com calma! Como posso afirmar que uma placa de rede que possui um endereço MAC ‘X’ não vai encontrar nenhuma outra placa com esse mesmo endereço? Especialmente se são várias empresas concorrentes fabricando placas de rede?”

Boa pergunta! O endereço MAC é composto por 48 bits, dos quais, os 24 iniciais representam a identificação do fabricante. Ou seja, duas placas de fabricantes diferentes já apresentam, de imediato, os conjuntos de 24 primeiros bits diferentes.

Ou seja, se duas placas são de fabricantes diferentes, elas já têm o início dos seus endereços MAC diferentes. E se duas placas são do mesmo fabricante, ele vai ter condições de controlar que não fará duas placas com o mesmo final.

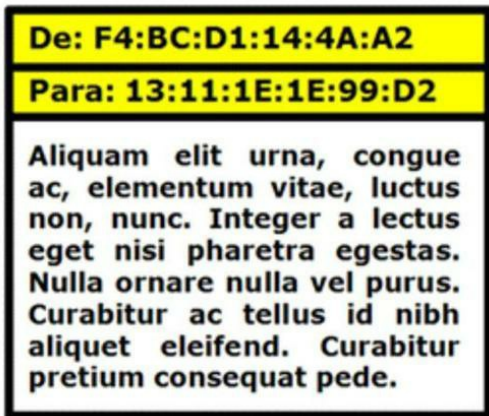
#### **8.8.1.2. Dando outra olhada na comunicação na rede**

Agora que conhecemos o endereço MAC e a sua função de identificar a rede, podemos esclarecer um pouco mais as comunicações nas LAN, que já vimos anteriormente nas arquiteturas de redes.

Em primeiro lugar, a ideia de pacote (quadro) agora parece mais clara se falarmos em quadro saindo de um computador, tendo como endereço de origem o MAC do micro remetente, e chegando a um computador de destino porque apresentava o endereço deste no campo de destino daquele pacote.

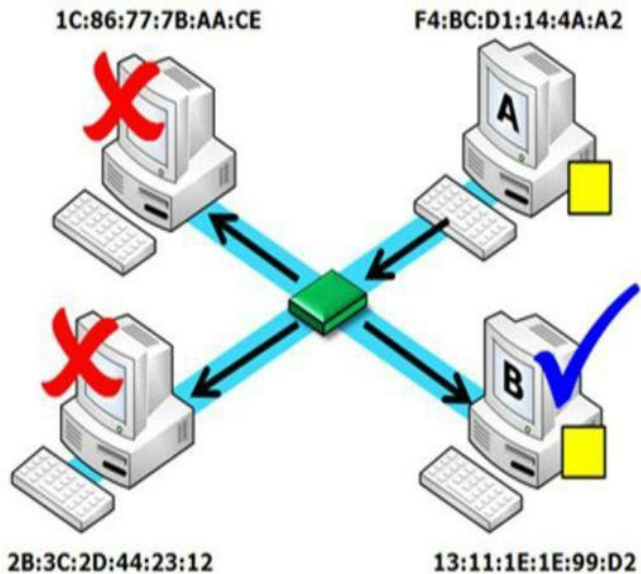
Ou seja, não importando a tecnologia ou o método de acesso da rede em questão, quando um computador envia um quadro à rede, ele o constrói, em uma LAN, usando seu endereço MAC como origem e o endereço MAC do destinatário como destino. Veja a figura a seguir.





*Figura 8.44 – Exemplo de um quadro com endereços MAC.*

Veja a seguir uma imagem que representa uma rede Ethernet (CSMA/CD): o micro “A” transmitiu o seu pacote (quadro) objetivando o micro “B”. O pacote tem seu cabeçalho preenchido com o endereço MAC do micro “A” no campo de origem e o endereço MAC do micro “B” como endereço de destino.



*Figura 8.45 – Funcionamento da Ethernet sob a óptica dos endereços MAC.*

O pacote é enviado e chegará a todas as placas de rede (por causa do broadcast). As placas desses micros vão ler o pacote (pelo menos o cabeçalho) e analisar o endereço de destino: caso o MAC contido no pacote seja igual ao da estação em questão (o que só acontecerá no micro “B”, no nosso exemplo), a mensagem será devidamente aceita.

Caso a placa de rede leia o pacote (pelo menos o cabeçalho) e note que o endereço MAC de destino descrito naquele pacote não é idêntico ao seu próprio MAC, a placa “se manca” e aceita que não é a destinatária real, descartando o pacote em questão.

Apesar de estar “adiantando” as coisas, gostaria de lembrar que a placa de rede é um equipamento pertencente à camada 2 (Camada de Enlace de Dados) do modelo OSI. (Muita calma nessa hora – veremos o que isso significa mais adiante.)

Outro lembrete: em qualquer comunicação entre componentes desta “tal” camada 2, usa-se o

termo **Quadro** em vez de Pacote (explicarei adiante) – por isso eu o utilizei aqui muitas vezes!

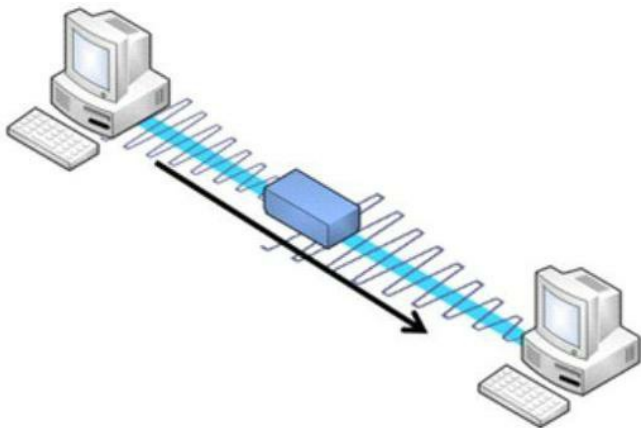
### 8.8.2. Repetidor

É um equipamento usado para regenerar o sinal elétrico (ou mesmo o luminoso) para que este possa ser transportado por uma distância maior.

Sabemos que os cabos usados nas conexões de rede convencionais possuem uma limitação de distância (cada tipo de cabo tem a sua), o que causa a atenuação (enfraquecimento) do sinal. Por isso, usamos repetidores para regenerar (gerar novamente) o sinal que se perderia pelo cabo.

Há repetidores para qualquer tipo de rede, mesmo para aquelas que não usam fios e, para essas, é apenas um ponto com antenas que retransmitem o sinal recebido.

Atualmente, não é muito comum encontrar um equipamento repetidor (apenas repetidor) no mercado. O mais comum é encontrar equipamentos diversos que acumulam a função de repetidores (como os hubs e switches atuais, que também servem como repetidores, regenerando os sinais que por ele passam).



*Figura 8.46 – Repetidor hipotético atuando.*

O repetidor é um equipamento que pertence à camada 1 (chamada de camada física) do modelo OSI. (Calma novamente... Já veremos isso!)

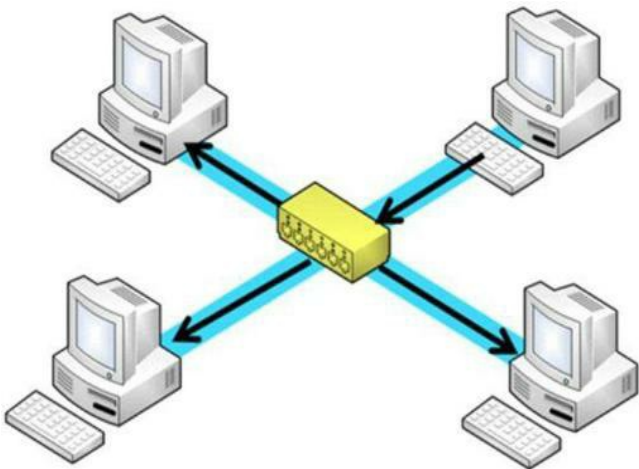
### 8.8.3. Hub

Um hub é um equipamento que serve como “centro” de uma rede Ethernet. Um hub é um equipamento simplório, que recebe os fios vindos dos micros (cabos de par trançado) e os conecta (conectores RJ-45) em sua estrutura. (Observe as diversas “portas” do hub.)



*Figura 8.47 – Um típico hub de cinco portas + porta Uplink*

Internamente o hub é apenas um barramento (uma conexão em topologia barra), o que explica seu funcionamento limitado e pouco inteligente. (Ele só funciona através de broadcast – ou seja, transmitindo para todos os demais micros). Provavelmente isso será perguntado desta maneira e você não pode errar: o hub Ethernet (hub comum) não faz nenhum tipo de filtro ou seleção sobre os dados que passam por ele. O hub sequer entende o que passa por ele. Os dados que são transmitidos passam pelo hub e, então, são imediatamente enviados a todos os demais computadores.



*Figura 8.48 – Um hub funcionando: necessariamente broadcast.*

“Então, se cair na prova que o hub sempre trabalha por broadcast, isso é CERTO?”

Sim! O hub não tem como trabalhar de outra forma, a não ser por broadcast, porque, internamente, ele é só um barramento (fios). Esse barramento conduz os sinais elétricos para todas as demais estações (porque, caro leitor, é isso que um “fio” faz, não é?).

Vamos estudar alguns tipos de hubs.

#### **8.8.3.1. Hub passivo**

Alguns hubs não precisam ser ligados à tomada elétrica, pois funcionam apenas como “conectores” para os fios. Estes são chamados hubs passivos.

*Um hub passivo não repete o sinal* (não atua como repetidor); portanto, o sinal que o atravessa vai perdendo sua força gradativamente sem a devida regeneração (atenuação). Um hub passivo é pouco usado hoje em dia.

#### **8.8.3.2. Hub ativo**

Os hubs Ethernet fabricados e vendidos atualmente são quase todos ativos. Um hub ativo é um

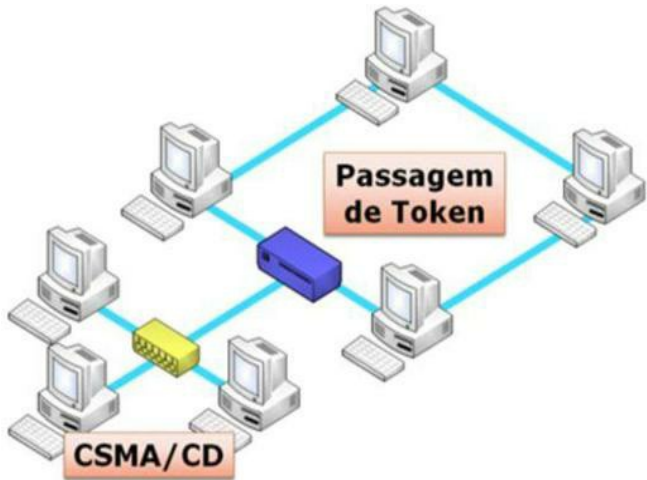
hub que se liga à tomada elétrica para repetir (regenerar) o sinal que o atravessa. Então, além de servir como um ponto de convergência de todos os cabos da rede, o *hub ativo atua como repetidor* para aumentar a potência do sinal, de modo que ele atravesse outros cabos de rede.

#### 8.8.4. Ponte

É um equipamento criado, originalmente, para interligar segmentos de rede de arquiteturas diferentes e permitir que eles se comuniquem normalmente. A ponte (bridge) é instalada entre um segmento de rede Ethernet e um segmento de rede Token Ring, por exemplo, e permite que os quadros (quadros de dados) passem de uma para a outra, caso seja necessário.

Devido à heterogeneidade de algumas redes locais, que podem apresentar variadas arquiteturas, como pedaços que usam Ethernet e outros que usam Token Ring, por exemplo, é necessário ligar esses “pedaços” para que se comuniquem. Mas há um “empecilho” para essa “união”.

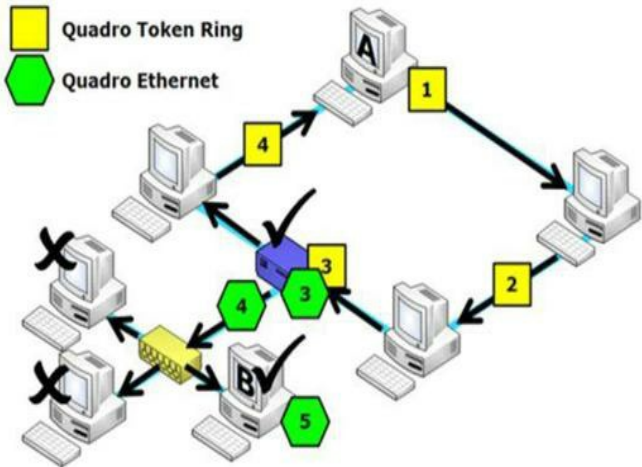
Tomando o exemplo anterior, em que analisamos uma rede formada por uma parte dos computadores ligados a um segmento Ethernet e os demais ligados a um anel na rede Token Ring, a ligação direta entre esses dois segmentos “mutuamente estrangeiros” não é possível. Na parte Ethernet, fala-se CSMA/CD e na parte Token Ring, respeita-se a Passagem de Token. Regras diferentes, protocolos de acesso diferentes. Em suma, linguagens diferentes. Esses dois segmentos não conseguem se comunicar diretamente sem o intermédio de um “intérprete”.



*Figura 8.49 – A ponte colocada como uma das estações do anel, ligando-o ao segmento Ethernet.*

A ponte servirá como tradutora dos quadros Ethernet, por exemplo, para quadros Token Ring e vice-versa. Isso permite que os quadros no formato Ethernet sejam convertidos em quadros que podem ser entendidos e retransmitidos na rede Token Ring.

Como a ponte funciona? Simples! Vamos a uma pequena história sobre a transmissão de um quadro do Micro “A” para o Micro “B” na rede mostrada na figura a seguir. A estrutura da rede é simples: um segmento Token Ring (onde o micro “A” se encontra) e um segmento Ethernet (onde o micro “B” está).



*Figura 8.50 – O histórico do quadro desde o micro “A” até o micro “B”.*

1. O micro “A” envia seu quadro, para que ele chegue ao micro “B”. Como já sabemos, o micro “A” faz isso colocando, no cabeçalho do quadro, o endereço MAC de “B” como destinatário.
2. Como é uma rede Token Ring, o quadro passa por todas as estações (que comparam o MAC localizado no destino do quadro com os seus próprios MAC). Como no exemplo anterior o MAC de destino é do micro “B”, as estações rejeitam a mensagem, passando-a adiante.
3. O quadro chega à ponte. A ponte, por sua vez, lê o quadro, identifica o endereço MAC de destino e, concluindo que este MAC é de uma estação do “outro lado”, ela traduz o quadro. (Na verdade, a ponte escreve outro quadro, desta vez no formato usado na arquitetura Ethernet, com o mesmo conteúdo do quadro Token Ring originário.) Observe a mudança de formato do quadro (na figura, representado pelo retângulo que vira um hexágono).
4. O quadro Ethernet é enviado (em broadcast) para o meio físico. (Para isso, a ponte precisou “escutar a portadora” a fim de saber se podia transmitir.) Enquanto isso, o quadro Token Ring original continua sua jornada até atingir o emissor da mensagem (o micro “A”).
5. Na rede Ethernet, já sabemos como funciona: todos os micros recebem os quadros,



analisam o endereço MAC contido no quadro e o comparam com seus próprios MAC. O micro destinatário (no caso, “B”) vai ver que seu MAC é idêntico ao MAC de destino existente no quadro e, então, vai aceitar o quadro e processá-lo. Os demais micros rejeitarão o quadro porque ele (o quadro) possui um endereço MAC diferente dos seus.

“Ei, João, quer dizer que a ponte sabe ‘ler’ o quadro, como a placa de rede?”

Sim, precisamente! A ponte é um equipamento da camada 2 (camada de enlace) do modelo OSI, assim como a placa de rede, por isso consegue ler quadros e trabalhar com endereços MAC.

Lembre-se também de que uma ponte pode ser usada, em alguns casos, para ligar dois segmentos de rede de mesma arquitetura (especialmente Ethernet).

“Peraí, João! Aí já não dá para aceitar! Se uma ponte é um dispositivo tradutor, para que serviria ligá-la entre dois segmentos que ‘falam a mesma língua’?”

Bom, sabendo que a ponte lê os quadros, ela serviria muito bem como um “ponto de segregação” da rede, uma espécie de “Muro de Berlim”.

“Como assim?”

Simples! Se uma ponte for colocada em um ponto estratégico da rede, ela consegue analisar quais quadros devem passar por ela (para o outro lado) e quais não devem.

Com esse tipo de filtro, quadros vindos de um setor de uma empresa, por exemplo, e endereçados para aquele mesmo setor não atravessariam toda a rede, mas seriam “bloqueados” pela ponte que saberia que eles não deviam passar.

O uso da ponte ligando partes de uma mesma arquitetura de rede, portanto, a transforma num dispositivo segmentador, mas não requer nenhum uso de sua função tradutora. Com a rede Ethernet dividida em segmentos bem definidos pela ponte, o número de colisões na rede diminui bruscamente, visto que agora o broadcast não atingirá necessariamente toda a rede.

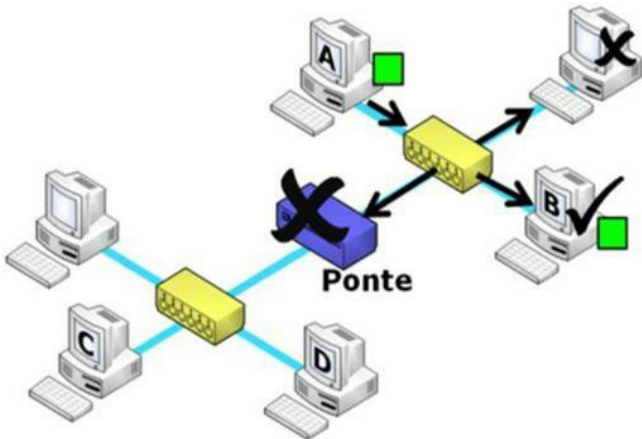


Figura 8.51 – Um quadro de “A” para “B” não precisa passar pela ponte – e ela sabe disso!

Quando uma ponte é colocada em uma rede Ethernet para separar a rede, chamamos cada “parte” resultante de *Segmento de Rede*, ou *Domínio de Colisão*. Portanto, a Figura 8.51 mostra uma rede com dois domínios de colisão (ou segmentos).

Um domínio de colisão é, portanto, uma área da rede de computadores onde quadros (ou pacotes) colidem, se duas estações tentarem acesso ao meio simultaneamente.

“Ei, João, essa definição não é meio ‘furada’? Se um micro ‘C’ no outro segmento tentar transmitir um quadro ao mesmo tempo em que o micro ‘A’ fizer a tentativa vai haver colisão, não vai?”

Não necessariamente, caro leitor! Se o micro “A” mandar um quadro para o micro “B” (eles estão no mesmo segmento, que chamaremos de segmento 1) e o micro “C” mandar um quadro para o micro “D” (ambos no outro segmento – o segmento 2), os dois quadros serão transmitidos perfeitamente (e ao mesmo tempo) porque eles não irão colidir.

O quadro enviado por “A” não passará para o segmento 2 (porque a ponte o cortará) e o quadro transmitido por “C” não passará para o segmento 1 (pelo mesmo motivo). Pronto, por causa da atuação da ponte, menos uma colisão no mundo!

#### 8.8.5. Switch

Nada mais é que um equipamento externamente semelhante a um hub (várias conexões para vários micros), mas que internamente possui a capacidade de chaveamento ou comutação (switching), ou seja, consegue enviar um pacote (um quadro, mais precisamente) exatamente para o segmento de destino.

Cada cabo (e micro) ligado ao switch está, necessariamente, em um segmento diferente, e não em um único barramento, como acontece no caso do hub.

Em outras palavras, o switch divide a rede em diversos segmentos, mais ou menos como a ponte. (A ponte só faz a segmentação da rede Ethernet em dois segmentos.) Além disso, a ponte faz o seu serviço por meio de software (programa) e o switch realiza essa segmentação diretamente no hardware (seus circuitos foram construídos para isso).



*Figura 8.52 – Switch com vários cabos ligados a ele.*

Devido às capacidades de chaveamento do switch, seu uso em uma rede Ethernet faz as colisões diminuírem bastante (em matéria de quantidade).

Há diversos switches para várias tecnologias de redes de computadores diferentes, como Ethernet, ATM entre outras. Vamos focar, claro, nos switches Ethernet, que são os mais comuns atualmente. (Devido ao fato de que essa tecnologia é a mais usada nos nossos dias.)

O switch, como já foi dito, tem condições de ler os quadros que por ele trafegam. Essa leitura é possível porque o switch possui processador e memória para realizar tais operações (ou seja, ele não é somente “uma caixa com um conjunto de fios” como o hub).

“Ô, João, por que o switch iria querer ler o quadro? É para descobrir o endereço MAC de destino escrito no cabeçalho daquele quadro?”

Precisamente. Complete seu raciocínio...

“E depois de ler o endereço MAC, o switch é capaz de enviar aquele quadro exatamente para o segmento em que o micro cujo MAC é igual àquele está localizado, não é?”

Perfeito! Em suma, o switch lê o quadro e, identificando o endereço MAC do destino, envia o quadro para o segmento exato.

“Mas, para isso, é necessário que o switch saiba previamente os endereços MAC dos micros ligados a ele, não é mesmo?”

Sim, sem dúvida. É para isso que existe a...

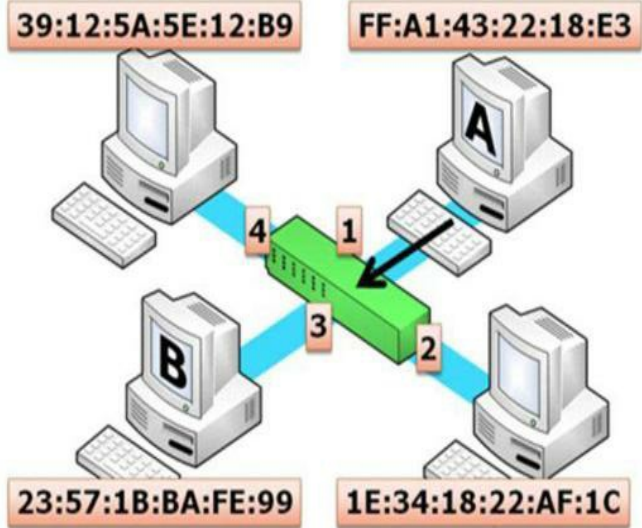
#### **8.8.5.1. Tabela de endereços MAC**

A grande maioria dos switches possui uma pequena quantidade de memória RAM (memória volátil e totalmente alterável, como já vimos) que tem como função, no processo de utilização do equipamento na rede, armazenar os endereços MAC das estações ligadas ao switch, permitindo que ele (o switch) possa realizar o seu trabalho de maneira mais eficiente.

“Sim, João, mas como funciona exatamente?”

É muito fácil. Imagine, observando a figura a seguir, que o micro “A” (cujo MAC é FF:A1:43:22:18:E3) transmite um quadro que chega ao switch. Supondo que esse quadro está endereçado ao micro “B”, cujo MAC é 23:57:1B:BA:FE:99, o switch será o responsável por enviar a mensagem unicamente ao micro “B”. Note, também, que o micro “A” está ligado à porta 1 do switch, enquanto o micro “B”, que será o destinatário, está ligado à porta 3.

Lembre-se de que o trabalho do switch é basicamente enviar o quadro ao micro “B”, ou seja, não mandar os sinais elétricos referentes àquele quadro para outra porta, além da porta 3, que é onde está conectado o micro de destino.



*Figura 8.53 – Micro “A” envia um quadro ao switch.*

Mas aí é que está a “alma” da sua pergunta, nobre leitor: o switch sabe ler o quadro; portanto, consegue, ao ler o cabeçalho daquele quadro, identificar o micro de destino. Mas como saber que o micro de destino (o micro com aquele endereço MAC) está ligado à porta 3 do switch?

O switch consulta a tabela que existe nele e vê que nela existe um registro indicando que o micro “B” (cujo MAC é 23:57:1B:BA:FE:99) está ligado à porta 3!

Mas, antes que você pergunte, e se o switch acabou de ser ligado? Se essa tabela está em uma memória RAM, quando o switch é ligado, ela inicia os trabalhos totalmente vazia. Mas vai “aprendendo” enquanto os quadros são transferidos pelo switch.

Um exemplo simples é o da figura anterior: na hora em que o quadro enviado pelo micro “A” passa pelo switch, ele (o switch) simplesmente lê o quadro (procurando o endereço MAC da origem no cabeçalho do quadro) e registra que o micro “A” (cujo MAC é FF:A1:43:22:18:E3) está ligado à porta 1.

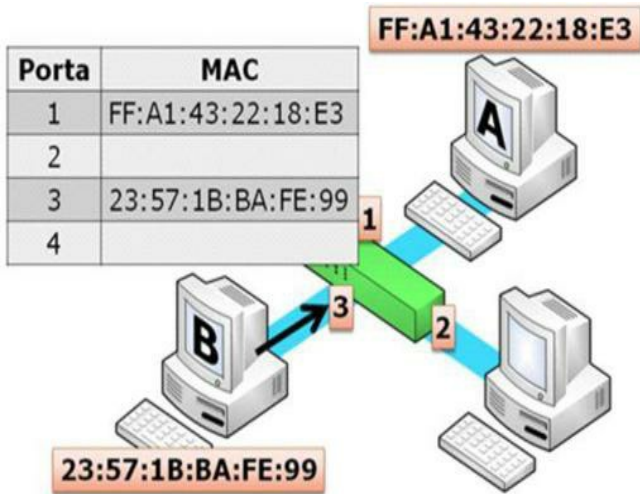


Figura 8.54 – Tabela de endereços MAC presente na memória do switch.

E, claro, quando o micro “B” responder (porque somente ele vai responder) e o quadro for transmitido pelo switch, este vai registrar que o micro “B” está ligado à porta 3. Com isso a tabela de endereços MAC será preenchida e se torna cada vez mais exata e precisa.

“OK, João, perfeito. Mas você ainda não respondeu... E se o switch não souber em que porta está o micro de destino? Quero dizer, e se aquele quadro for o primeiro a ser mandado para o micro ‘B’ e a tabela não possui registro sobre a localização dele?”

Ah... Claro! Desculpe! Quando o switch não sabe quem é o micro exato de destino (ou seja, quando ele não sabe em que porta o micro de destino está conectado), ele faz uso de um recurso “arcaico, mas eficaz”: o **broadcast**! Simplesmente o switch manda a mensagem “conscientemente” para todas as portas. Quando a resposta for dada, e o quadro de resposta passar pelo switch, ele já terá condições de registrar o MAC daquele micro na sua tabela.

Portanto, o switch pode usar broadcast. (Ele só usa quando precisa.)

Por ter condições de “compreender” quadros para ler endereços MAC neles contidos, o switch é classificado como um dispositivo da camada 2 (camada de enlace), assim como a placa de

rede e a ponte. Novamente, tenha calma. Veremos isso adiante. (Sei que já deve estar “enchendo o saco”!)

#### 8.8.6. Ponto de acesso (Access Point)

Como já foi visto rapidamente, para que uma rede de computadores Wi-Fi seja montada em modo conhecido como infraestrutura, é necessária a presença de um equipamento que centraliza todas as comunicações desta rede. Esse equipamento é conhecido como ponto de acesso Wi-Fi ou simplesmente ponto de acesso. (Alguns livros não traduzem o termo do inglês, portanto se referem a ele como AP – Access Point.)



*Figura 8.55 – Ponto de Acesso Wi-Fi (já com MIMO) da USRobotics®.*

Cabe ao ponto de acesso (e das placas de rede Wi-Fi) tratar de questões como evitar as colisões (CSMA/CA), criptografar e descriptografar os quadros que se encontram em redes que usam segurança (WEP ou WPA), entre outras tarefas.

O ponto de acesso é, assim como ponte, placa de rede e switch, um equipamento da camada 2 (camada de enlace).

#### 8.8.7. Roteador

Roteador (ou router) é o nome dado a um equipamento capaz de rotear! Rotear significa definir a rota. Um roteador é um equipamento que, em suma, define a rota a ser percorrida pelos pacotes da origem ao destino.

“Ô, João, você não quis dizer ‘quadros’ em vez de ‘pacotes’?”

Não. No caso do roteador, devemos chamar de pacotes mesmo! O roteador é um equipamento descrito como pertencente à camada 3 (camada de redes) – ou seja, ele é mais “especializado” que o switch, a ponte e o ponto de acesso.

“Tudo bem, mas em que consiste essa ‘especialização’? No que ele se diferencia dos equipamentos já vistos?”

É simples, mas preste atenção, caro leitor! O roteador não serve para interligar computadores ou segmentos dentro de uma mesma rede. O roteador serve para *interligar redes distintas*! Ou seja, ele não liga dois ou três micros em uma rede; liga duas ou três redes em uma estrutura conhecida como inter-redes (ou Inter-net).



*Figura 8.56 – Um roteador da Cisco®.*

A figura a seguir mostra um exemplo de Inter-net (ou Inter-Networking, que traduzindo seria “estrutura de ligação entre redes”).



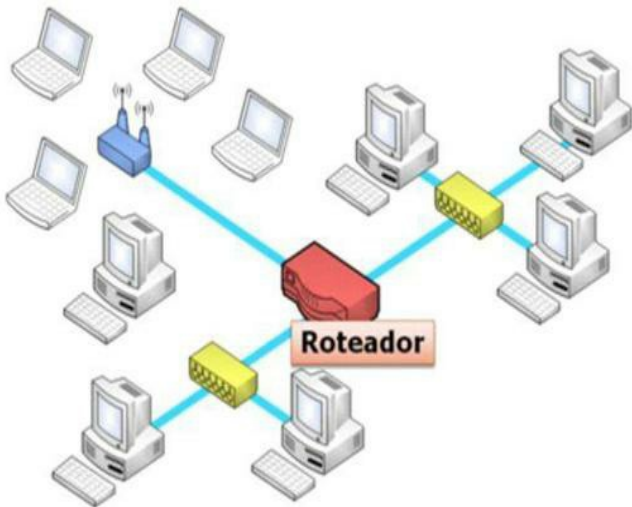


Figura 8.57 – Um roteador ligando três redes distintas. (Não são três segmentos. São três redes!)

Algo interessante aqui é: o endereço MAC não é o mais importante nas comunicações entre redes. O endereço MAC de cada placa de rede é imprescindível nas comunicações que se processam em uma única rede. (Quando uma placa de rede quer se comunicar com outra na mesma rede.) Em redes diferentes, surge uma nova forma de localização e identificação de origem e destino: o **endereço lógico**.

É o seguinte: o endereço MAC é chamado de endereço físico, pois está contido em cada placa de rede em sua memória ROM. Esse endereço é usado nas comunicações que acontecem dentro de uma única rede (sem ter de passar pelo roteador). Mas, quando há necessidade de comunicação com computadores em outras redes (ou seja, a mensagem tem de passar pelo roteador da rede), o endereço MAC perde, em muito, a sua importância, pois o roteador lê, a priori, um endereço de maior abrangência, chamado de endereço lógico (que, na Internet, é chamado de **endereço IP**).

“Quer dizer que o roteador não pode ler endereços MAC? Só consegue ler esses tais endereços IP?”

Não! O roteador lê endereços MAC, pois ele vai precisar disso para enviar os pacotes na forma de quadros na rede de destino. Veremos isso com mais detalhes depois que aprendermos sobre endereços IP (depois dos modelos de camadas).

A questão do roteador é que, para o desempenho de sua função, o endereço IP é mais importante que o endereço MAC. E é conhecendo o endereço IP do micro de destino que se descobre o seu endereço MAC.

“Hã?!”

Calma... Só veremos isso depois de vermos modelos de camadas e os protocolos do TCP/IP.

#### 8.8.7.1. Tabela de roteamento

Assim como os switches possuem uma tabela que associa os endereços MAC dos micros ligados a ele com as portas que ele possui, o roteador também traz uma tabela que associa os endereços IP dos micros e/ou outros roteadores às suas portas (interfaces, como chamamos). Essa tabela é chamada **tabela de roteamento**. Ela traz informações das faixas de endereços IP das redes ligadas àquele roteador, ou seja, cada roteador sabe a quais redes ele está ligado.

A tabela de roteamento pode ser estática (definida manualmente pelo administrador da rede – nesse caso, é claro, os roteadores são mais baratos) ou dinâmica (em que o roteador “aprende”, com o uso, os endereços dos outros roteadores e das outras redes a ele ligadas, além de receber informações das tabelas de outros roteadores).

Portanto, os roteadores trocam informações sobre suas tabelas durante o uso da rede, para tornar mais rápida e precisa a comunicação. Ou seja, os routers “fofocam”.

Lembre-se: ao ambiente formado por várias redes distintas interligadas, chamamos de **inter-redes** (ou **inter-net**, de onde derivou a famosa Internet); portanto, se alguma questão perguntar: “qual é o equipamento que interliga redes distintas?”, sabemos que é o roteador.

#### 8.8.8. Vários componentes de rede juntos

É muito comum encontrar atualmente equipamentos de rede que “acumulam” funções de diversos componentes, como ponto de acesso, switch, modem de banda larga (ADSL ou cabo) e roteador num só corpo físico.

Nesse caso, não podemos simplesmente analisar como um único dispositivo (apesar de comprá-lo como tal). Continuam valendo os conhecimentos que adquirimos quando analisamos cada um deles. Veja um exemplo na figura a seguir.



*Figura 8.58 – Ponto de acesso + switch + roteador + modem ADSL em um só produto.*

Então, ser “roteador” hoje em dia é basicamente acumular a função de fazer o roteamento, porque encontrar o roteador como equipamento separado é, digamos, um tanto incomum (pelo menos, para nós, usuários domésticos) – embora seja perfeitamente possível.

### 8.9. Modelos de camadas

Eis aqui um assunto que só a Esaf (pelo menos até agora) tem coragem de exigir em provas de concursos para qualquer área (fora da área específica de informática). Ah... A Fundação Getúlio Vargas (FGV) também! Vamos a ele.

Já discutimos que cada tecnologia de rede de computadores – antiga, atual ou futura – deve ser padronizada, ou seja, deve ter seu funcionamento descrito em um documento conhecido como padrão (assinado por algum órgão competente) para que se torne comercialmente viável.

No caso de você querer criar uma nova tecnologia de rede e descrever o funcionamento desse seu novo “ambiente” para os órgãos competentes poderem padronizá-la e para que ela se torne comercialmente utilizável, por onde você começaria? Além disso, como a “sua nova rede” poderia ser descrita de modo que fosse comparada facilmente às tecnologias já existentes, garantindo, inclusive, que ela se comunique com as demais?

Sua rede (e qualquer uma) deve ser descrita (desenhada ou projetada, se você preferir) de uma forma curiosa: através de uma técnica conhecida como **modelo de camadas**. Modelo de camadas é, em suma, uma maneira de “explicar uma rede”, uma forma de “entender” a rede. Ou ainda, um modelo de camada é uma metodologia para simplesmente “visualizar” como uma

rede funciona.

Vamos lá... Esse é um assunto muito teórico e chato para quem não é do ramo (e, às vezes, até para quem é da área de Informática), mas vou tentar torná-lo o mais palpável possível, OK?

Como você explicaria a alguém como “fritar um ovo”? Vamos lá. Tente...

“Eita, João, perai... Que tal: coloque um pouco de manteiga na frigideira, ligue o fogo, quebre o ovo e jogue clara e gema na frigideira enquanto mexe-os junto à manteiga até que chegue à consistência desejada. Tá bom assim?”

Bem, se eu fosse depender de você, caro leitor, para me ensinar a cozinhar, eu estaria morto de fome!

“Tá bom, João, não precisa zoar!”

Estou brincando. Vamos lá... Você fez o melhor que pôde, mas se eu pedisse para me explicar como fazer um bolo, a “forma de explicar” mudaria tanto quanto os ingredientes, não é? Quero dizer, não há uma normatização, uma padronização na sua forma de explicar.

Vamos propor o seguinte: ao explicar sobre qualquer tipo de prato gastronômico, que tal fazê-lo preenchendo um “gabarito”, como um “formulário de inscrição” de um concurso: esse formulário solicita as seguintes informações:

|   |                |
|---|----------------|
| 5 | Aplicação:     |
| 4 | Guarnição:     |
| 3 | Modo de fazer: |
| 2 | Ingredientes:  |
| 1 | Utensílios:    |

Para preencher corretamente esse formulário, deve-se entender o que deverá estar em cada um dos campos (ou camadas, como vamos chamar a partir de agora).

A camada 1 (utensílios) apresentará descrições dos vários apetrechos e eletrodomésticos usados na criação do seu prato. Ou seja, itens como espátula, colher de pau, fogão, freezer, wok, garfo para carne, entre outros, pertencem a essa camada.

A camada 2 (ingredientes) tratará de descrever as “matérias-primas” usadas para dar vida ao seu empreendimento gastronômico. Isso significa que salsa, cebolinha, carne, filé, salmão, couve, cenoura, sazón (esse é importantíssimo: tudo o que faço na cozinha depende dele!), arroz, shoyu etc. devem ser colocados nessa camada.

A camada 3 (modo de fazer) descreverá os processos a serem realizados na criação do prato,

como “coloque na chapa até dourar”, “não use o micro-ondas”, “misture o vinagre ao arroz frio e vá mexendo até obter uma liga”. Todas essas “regras” fazem parte da camada 3 do nosso modelo, pois são componentes do modo de fazer dos pratos.

A camada 4 (guarnição) descreve os acompanhamentos do prato como enfeite (para os pratos mais “chiques”). Ou seja, fatias de laranja, tomate-cereja cortado em formato de rosa, um morango grande, folhas de louro são exemplos de itens que merecem ser incluídos na camada 4.

Finalmente, a camada 5 (aplicação) descreve simplesmente em que ocasião o prato se aplica, como: prato principal à noite, sobremesa, almoço, entrada etc. devem ser colocados nessa camada, quando existirem.

Deu para entender que cada camada existe para que os componentes certos estejam exatamente nelas, não é? Por exemplo, um fogão não poderia pertencer à camada de ingredientes, bem como “untar toda a assadeira” não pode ser considerado componente da ocasião.

Note que, ao definir isso, todo processo que você já conhecia para cozinhar pode ser sistematizado seguindo esse modelo. Então, fritar o ovo pode ser explicado assim:

Camada 1: frigideira, garfo, espátula, fogão.

Camada 2: manteiga, ovo, sazón (segredo de família).

Camada 3: coloque a manteiga na frigideira no fogo médio; jogue uma pitada de sazón (sabor de “ervas” fica delicioso) e mexa um pouco; quebre um ovo grande; mexa até a consistência certa (prefiro “ovos mexidos”).

Camada 4: pão, tomates, cebolas, salsa, torradas.

Camada 5: café da manhã, jantar.

Ou seja, fazer um bolo, fritar um ovo, cozinhar o feijão, assar uma pizza, fazer sushi podem ser descritos usando o modelo de camadas mostrado anteriormente. Isso significa que você acabou de aprender, tecnicamente, a cozinhar qualquer coisa! (Livro de informática + livro de culinária pelo preço de um? Que promoção, hein?)

Então, você acabou de ver a “explicação” do “ovo frito” em um modelo de camadas que criamos. Note que se quiséssemos criar um modelo de camadas com 10 camadas, ou com apenas 3 camadas, só dependeria de nós. O modelo de camadas deve ser adequado à realidade do objeto de estudos dele.

“Certo, João, mas se a gente for fazer outro prato, como um bolo, os componentes mudam!”

Sim, é claro! É outro prato. Mas o MODELO não muda. Ou seja, o segredo é que o modelo de camadas é único para todos os pratos. Assim como nas redes de computadores (que são variadas), o modelo de camadas é o mesmo e isso garante a compatibilidade entre os vários tipos de redes de computadores.

Final, entenda isto: não importa qual o prato a ser feito, é necessário que haja utensílios (sim, mesmo que sejam diferentes para cada prato). E os ingredientes? Claro que haverá ingredientes em todos os pratos (novamente, mesmo que sejam diferentes de um prato para outro.) – o conteúdo de cada camada pode mudar, mas se esquematizarmos o modelo com aquelas camadas, elas nunca mudam. As camadas só mudam se mudarmos o modelo de camadas, ou seja, se mudarmos a forma de explicar.

Será que ficou claro? Espero que sim...

Quanto às redes de computadores, há um modelo de camadas que é seguido (muitas vezes

com algumas alterações) por todos que criam e modificam tecnologias de rede, o modelo de camadas OSI (Open Systems Interconnection – Interconexão de Sistemas Abertos), desenvolvido pela ISO (International Standardization Organization – Organização Internacional de Padronização). Esse modelo é normalmente conhecido como **Modelo de Camadas ISO/OSI**.

### 8.9.1. Modelo de camadas ISO/OSI

O modelo OSI é composto por sete camadas diferentes e é um marco da padronização de redes de computadores. Na prática, ele não é seguido à risca pelas empresas que atualmente trabalham com tecnologias de redes, mas é a partir desse modelo que novos modelos são criados. Por esse motivo, o modelo OSI é chamado **Modelo de Referência ISO/OSI**.

Além do modelo OSI, o modelo TCP/IP também será estudado neste livro, devido à sua ampla utilização na Internet e, por isso, será o mais exigido em provas de concursos. O modelo OSI é pouco exigido em provas abertas, mas a Esaf já o exigiu em provas de Auditor-Fiscal da Receita; portanto, se é seu objetivo, leia todo este tópico.

As sete camadas do modelo de redes OSI são:

#### 8.9.1.1. Camada 1 – camada física

Descreve os equipamentos físicos usados na transmissão dos sinais brutos (elétricos, luminosos ou eletromagnéticos) e os meios de transmissão. São integrantes desta camada os cabos (UTP, fibra óptica, coaxial), os repetidores, os conectores (RJ-45, BNC), as ondas de RF, as ondas infravermelhas e os hubs.

A preocupação desta camada não é com o significado dos dados transmitidos (pacotes, mensagens), mas sim com a forma física de sua transmissão (voltagem correta para determinar os bits 0 e 1 elétricos, corrente elétrica, frequência de transmissão, duração do bit), ou seja, a forma “bruta” dos sinais que transmitem dados.

“Ah, João, agora entendi por que o hub está na camada 1. Quando eu voltar a ler o tópico que fala nos equipamentos de rede, tudo será mais claro!”

Exatamente! O hub pertence à camada 1 porque ele não consegue entender nada além de sinais elétricos. Nesse sentido, o hub é tão “inteligente” quanto um fio, afinal, ele é somente um conjunto de fios mesmo!

Lembre-se: todo componente que não consegue ler as informações que passam por ele (nem como quadros, nem como pacotes, nem como mensagens), ou seja, entende tais informações apenas como elas realmente são: pulsos elétricos, luminosos ou eletromagnéticos e nada mais, é um componente localizado na camada 1.

Para terminar: se qualquer equipamento da camada 1 (hubs, fios, repetidores) pudesse “dizer o que está vendo”, diria que por ele estão passando somente vários 001001010100101010101.

#### 8.9.1.2. Camada 2 – camada de enlace (ou enlace de dados)

Esta camada é responsável por “reunir” os sinais brutos (zeros e uns) e “entendê-los” como quadros, identificando suas origens e destinos (endereços MAC) e corrigindo possíveis erros ocorridos durante a transmissão pelos meios físicos.

Como os dispositivos da camada 1 são apenas “fios” (ou seja, transmitem sinais brutos, sem

nenhum grau de “inteligência”), torna-se responsabilidade dos dispositivos da camada 2 detectarem (e, se possível, corrigirem) as besteiras que a camada 1 venha a cometer.

Como vemos, qualquer dispositivo que consiga entender os quadros e ler os endereços MAC, permitindo, assim, a comunicação dentro de uma única rede (ou seja, qualquer comunicação que não “atravesse” um roteador) está automaticamente classificado como pertencente à camada 2. Os equipamentos físicos que merecem pertencer à camada 2 são a placa de rede, a ponte, o ponto de acesso e o switch.

Os protocolos CSMA/CD, CSMA/CA e as diversas tecnologias de rede (Ethernet, FDDI, Token Ring, ATM, IEEE 802.11 etc.) também são descritos como pertencentes a essa camada, pois regulam, justamente, a comunicação entre computadores dentro de uma única rede. Ou seja, na camada 2 não existem apenas componentes físicos, mas lógicos (protocolos) também.

“João, por que essa insistência em ‘numa única rede’?”

Simples, caro leitor! Quando uma comunicação é realizada entre duas redes diferentes, ela necessita de um equipamento roteador e passará a ser efetuada usando os endereços lógicos (que conheceremos como endereços IP). Isso já é responsabilidade da camada 3 (camada de rede), na qual estão inseridos o roteador e o endereço IP.

Lembre-se: qualquer equipamento que seja responsável por estabelecer, realizar e encerrar a comunicação entre duas estações dentro de uma mesma rede (um mesmo “enlace”), sendo capaz, para isso, de ler e interpretar os endereços MAC presentes nos quadros é, sem dúvidas, pertencente à camada 2 (camada de enlace).

E tem mais... Não são só equipamentos, não! Protocolos de acesso ao meio e arquiteturas de LANs, MANs e WANs também são considerados itens da camada 2!

Ahhh... Eu quase ia me esquecendo... Falou-se em quadros (frames), falou-se em camada de enlace (camada 2). Falou-se em endereços físicos (endereços que estão presentes nas próprias interfaces de rede – as placas de rede), como os endereços MAC, falou-se em camada 2 também!

### 8.9.1.3. Camada 3 – camada de rede

É a camada em que se localizam os equipamentos e protocolos responsáveis por interligar diversas redes. Os equipamentos (e protocolos) que criam e mantêm um ambiente inter-redes (inter-net), como o roteador, por exemplo, são pertencentes à camada 3. Vamos a um comparativo entre a camada 2 e a camada 3.

Quando a comunicação se dá dentro de uma única rede, como vimos, as estações (computadores) envolvidas reconhecem-se mutuamente pelos seus endereços MAC (endereços físicos). Toda a comunicação é feita por meio de pequenos pedaços de informação chamados quadros (frames) devidamente identificados com o endereço MAC da origem e o endereço MAC do destino.

Tudo isso é ambiente “de camada 2”. Todos os envolvidos (placas de rede, switches, pontes, CSMA/CD, passagem de token etc.) são pertencentes à camada 2.

Porém, quando a comunicação “extrapola” uma rede, “transborda” para outras redes (ou seja, quando uma mensagem tem de sair da rede em que o micro de origem está para chegar a outra rede, onde o destino se encontra), é necessário que “entrem em ação” equipamentos e

protocolos diferentes, capazes de “se virar” (ter “jogo de cintura”) para propiciar a comunicação nesse ambiente.

Para começo de conversa, em uma comunicação inter-redes (entre redes distintas), não se usa, a priori, o endereço MAC (endereço físico), mas outro endereço, válido para a estrutura de Internet inteira, chamado endereço IP (é o mais usado hoje).

Roteadores leem endereços IP e é por isso que podem encaminhar (rotear) um pacote entre uma rede e outra. Portanto, os roteadores (e todos os protocolos, como o IP) são descritos na camada 3 (camada de redes) porque possibilitam a comunicação entre redes distintas.

“Pacotes, João? O nome não seria ‘quadros’?”

Não, caro leitor. Na camada 3 chamamos de pacotes mesmo! Mas isso é meio “burocrático”. Vamos ver mais adiante a ideia de PDU e entender isso melhor!

“Então qual seria, afinal, a diferença entre um ‘quadro’ e um ‘pacote’?”

Ambos são “pedaços” de informação, caro leitor. A principal diferença é que um quadro pode ser transportado apenas por um único enlace físico (uma única rede), pois o endereço que dá identificação de origem e destino para os quadros (endereço MAC) só tem “competência” dentro de uma única rede. Um pacote pode ser enviado entre redes diferentes, porque usa, como identificador, um endereço que atua em um “cenário” mais abrangente, envolvendo diversas redes diferentes (esse endereço é o endereço IP – ou endereço lógico).

Então, resumindo, você tem de aceitar isto: falou-se em “camada 3”, então pense, imediatamente, em equipamentos e protocolos (regras) para a comunicação entre redes distintas. Falou em endereço IP, em vez de endereço MAC, é camada 3. Qualquer equipamento (roteador, por exemplo) que consiga ler endereços IP pertence à camada 3.

#### 8.9.1.4. Camada 4 – camada de transporte

Até agora, vimos camadas muito próximas ao hardware, ou seja, muito próximas à comunicação entre as máquinas, os roteadores (camada 3) e as placas em uma única rede (camada 2); ou ainda a comunicação pura no fio (camada 1).

Chegou a hora de analisar a comunicação sob outra óptica (e é essa a diferença entre as camadas – apenas a óptica sob a qual se veem as mensagens durante a comunicação – mas isso é assunto para mais adiante). Enfim, chegou a hora de analisar a comunicação sob a óptica da mensagem e não da troca de pacotes e/ou quadros.

Uma mensagem é qualquer bloco fechado de informações que se deseja transmitir, como um e-mail, um arquivo PDF, uma foto, uma página da Internet, uma música em MP3, qualquer arquivo que se deseja transmitir pela estrutura das redes. Já havíamos visto isso!

A questão é que, dependendo da estrutura das redes envolvidas, a mensagem (inteira) não pode ser transmitida sem que antes seja dividida em pequenos pedaços (os pacotes – quando analisados sob a óptica da camada 3 – ou quadros – quando vistos na camada 2).

E tem mais. Não importa qual o tamanho do pacote ou quais as características do quadro daquela tecnologia de rede em si, as mensagens originalmente escritas em bom português serão transmitidas, pelos fios, da maneira como sempre foram... 001001010010101001010101 (pulsos físicos que significam “0” e “1”).

Então, uma mensagem escrita em português segue um longo caminho desde o momento em



que é digitada pelo usuário até o momento em que começa a trafegar pelos fios. E a participação da camada 4 é vital para o funcionamento desses processos.

A camada 4 (camada de transporte) tem como responsabilidade oferecer meios de **controle da transmissão**: métodos e técnicas que permitam a perfeita conversa entre origem e destino, de modo que a mensagem inteira que saiu consiga chegar perfeitamente, mesmo que isso leve centenas de pacotes que passarão por dezenas de redes distintas.

“João, não deu para perceber a exata função da camada 4. Explica de outra maneira, por favor!”

Claro! Entenda que a mensagem (e-mail, por exemplo) é uma entidade única. Ninguém envia ½ e-mail. (Embora alguns digam que “2 mails = 1 inteiro” – hehehe – piada infame, desculpe... Só tem graça quando ela é “ouvida”, não “lida”...)

Voltando ao assunto: sempre enviamos um e-mail (inteiro), não importa seu tamanho ou a quantidade de anexos que ele possui (essa característica influencia diretamente no tamanho do e-mail). Então, se eu envio um e-mail, quero que esse e-mail, inteiro, chegue ao destino, não é mesmo?

Mas para que isso ocorra, esse e-mail tem de ser dividido em diversos pedaços para atravessar a Internet (pacotes) e, em cada rede por onde ele passar, deverá adequar-se à tecnologia daquela rede sendo colocado em quadros.

Pois bem. Como sabemos que a mensagem será dividida, é interessante ter uma camada que faça a divisão de maneira adequada (separe as mensagens) no emissor, atribua-lhe números de controle (como “pacote 1 de 15”, “pacote 2 de 15”, “pacote 3 de 15” etc.) e, quando estes chegarem ao micro de destino, a mesma camada naquele micro possa unir os pacotes enviados em ordem correta, resultando, assim, na montagem perfeita da mensagem original.

Essa camada “separadora”, “conferente” e “juntadora” é a camada de transporte. A camada 4 não se responsabiliza por mandar os pacotes de roteador em roteador, tampouco é responsabilidade dela mexer com os quadros, enviando-os entre as estações numa rede. Esses são trabalhos das camadas inferiores.

A camada de transporte é responsável pela **comunicação fim a fim** (origem-destino). Ela é responsável pela perfeita troca de mensagens entre o emissor (que as separa em pedaços, atribuindo uma ordem a eles) e o receptor (que recebe tais pedaços e os junta ordenadamente). É a camada de transporte que também detecta e corrige possíveis erros em pacotes. Também é a camada de transporte que detecta se algum pacote estiver faltando. (Isso pode acontecer.)

Na camada de transporte não existem equipamentos. (Quer dizer que historicamente não há equipamentos – dispositivos físicos – que mereçam ser classificados como pertencentes a essa camada.) Mas há protocolos. Os mais importantes são o TCP e o UDP (usados na Internet – isso quando analisamos o modelo de camadas usado na Internet – que não é o OSI), mas também há outros, como o SPX, para as redes Novell Netware.

#### 8.9.1.5. Camada 5 – camada de sessão

Esta camada não saiu do papel (hoje em dia, são usados modelos que não utilizam essa camada), mas, como ela está descrita no OSI, precisa ser explicada.

Segundo o modelo OSI, dois computadores que desejam se comunicar precisam, antes de

qualquer outra coisa, estabelecer, entre eles, um “acordo de transação”, ou seja, antes de transmitirem entre si o primeiro pacote, os micros envolvidos devem iniciar um “cenário”, um “ambiente”, um “momento” oficial de comunicação ininterrupta. Esse momento é chamado de sessão.

Fazendo uma comparação bem oportuna: quando vamos ligar (pelo telefone mesmo) para alguém conhecido a fim de contar uma novidade, é necessário que haja, antes da primeira palavra trocada entre os dois envolvidos, o estabelecimento da conexão telefônica, não é? Essa ligação é uma sessão telefônica que só será interrompida quando o “ligador” desligar seu telefone.

Sessão é, portanto, uma *relação ininterrupta de comunicação*. Uma transação. Um procedimento que tem início e fim.

A camada de sessão determina as regras e “burocracias” para o estabelecimento de tais sessões. Em suma, a função dessa camada é gerenciar o estabelecimento de sessões de comunicação. Ou seja, todas as regras, exigências e determinações presentes na camada de sessão (até mesmo quanto ao seu simples objetivo) são apenas “teoria”, já que ela nunca foi posta em prática.

Lembre-se disto, meu amigo leitor: a camada de sessão está descrita e especificada no modelo de camadas ISO/OSI, e somente aí! Portanto, como o modelo OSI é apenas teoria (não é usado, na prática, em lugar algum), a camada de sessão também é apenas teoria.

#### **8.9.1.6. Camada 6 – camada de apresentação**

Outra camada “cabeça de bacalhau”. Ninguém nunca viu!

Esta camada, assim como a de sessão, é descrita apenas no modelo OSI e em mais nenhum modelo de camadas prático! Portanto, é perfeitamente possível deduzir que essa camada também é “teoria” – apenas utopia.

A que se propõe a camada de apresentação (claro, se perguntarem na prova sobre ela)? Simples! Basicamente conversão! A camada de apresentação tem a árdua tarefa de se comunicar com a camada de aplicação (que está intimamente ligada aos usuários). Da camada de aplicação (camada 7), provêm os mais variados tipos de informação (e-mail, páginas, arquivos PDF, arquivos MP3) que precisam ser transformados (digamos “traduzidos”) para um formato geral, um formato que isentasse a camada de transporte de problemas para separar os pacotes.

Esse processo de transformar as mensagens de formato variado em um formato genérico único, que servirá para facilitar todo o processo de transmissão que se segue, inclui procedimentos como criptografia (reescrita embaralhada das informações) e compactação.

Portanto, não se esqueça disto: a camada de apresentação traduz as mensagens vindas da camada de aplicação para um formato genérico antes de serem transmitidas. Além disso, criptografia e compactação também são tarefas desempenhadas por essa camada.

#### **8.9.1.7. Camada 7 – camada de aplicação**

O mais alto nível da pilha OSI é a camada de aplicação, que entra em contato diretamente com o mundo exterior, ou seja, nós, os usuários.

Nessa camada são descritos protocolos que realizam diretamente as tarefas a que temos acesso, como e-mails, navegação na Web, transferência de arquivos, bate-papo etc. Esses protocolos são chamados protocolos de aplicação.

Os próprios serviços que podemos desempenhar (como o envio e recebimento de mensagens de e-mail e a navegação em páginas Web) são descritos como pertencentes a essa camada.

Então, é fácil lembrar: os protocolos e serviços (tarefas) a que os usuários têm acesso são componentes da camada de aplicação. Essa camada recebe a mensagem pura, escrita diretamente pelo usuário, e manda para as camadas mais baixas (claro que, diretamente, para a camada de apresentação).

#### **8.9.1.8. O modelo ISO/OSI completo**

A seguir, o desenho esquemático do modelo OSI completo:



*Figura 8.59 – Modelo de referência ISO/OSI.*

“Mas, João, afinal... Para que isso serve mesmo? Não é apenas ‘conversa para boi dormir’, não, né?”

Se dá para um boi dormir, eu não sei... Mas eu (que tenho quase o peso de um boi) dormi, sim, durante a explicação que alguns professores me fizeram sobre isso. Mas eu queria dizer que o único – se não for o único, é o mais importante – motivo para aprender isso é que vão cobrar isso em prova (especialmente a Esaf! Não se iluda! A Esaf, vez por outra, cobra “esse trem todo”!).

Em segundo lugar, o modelo de camadas é muito importante para as pessoas que estudam redes de computadores porque através dessa técnica de “sistematização”, essa “linha de raciocínio”, é possível entender de maneira bem simples e organizada como as redes funcionam.

#### **8.9.1.9. Exemplificando a comunicação (visão do modelo OSI)**

Como é para “sistematizar” ou “explicar” as comunicações entre computadores, os modelos de camadas podem ser aplicados para analisarmos, sob a óptica deles, qualquer tipo de troca de informações, como um inofensivo e ingênuo e-mail etc.

Note como sua “percepção” acerca do e-mail vai mudar um pouco. Antes de começarmos, porém, é necessário que você saiba que as camadas vão “modificando” a mensagem à medida que esta vai descendo por elas (sim, descendo!).

Essa modificação tem um objetivo: cada camada, no micro emissor, coloca “informações” ou “marcações” (como selos postais ou carimbos) para que a exata mesma camada, no micro receptor, tenha condições de entender e aceitar a mensagem (ou pacote, ou quadro) em questão.

Então, a ordem é a seguinte: a mensagem parte da camada de aplicação do usuário remetente (claro, pois parte de um ser humano que a escreveu); atravessa todo o modelo de camadas no emissor; parte efetivamente para o receptor pela camada física (fios, cabos, antenas); atravessa (subindo) todas as camadas no receptor; enfim, chega à camada de aplicação, onde será lida pelo usuário destinatário.

O caminho é simplesmente em “U” – de aplicação, na origem, para aplicação, no destino. Essa é a ordem através da qual as mensagens são transmitidas de um computador a outro, analisando-as segundo a ideia de modelos de camadas, claro! Veja a figura a seguir e entenderá mais facilmente.

# ORIGEM

# DESTINO



Figura 8.60 – Comunicação analisada pelo modelo OSI.

Observe, nessa figura, que há setas brancas perfazendo o caminho real (em “U”) e há setas escuras indicando a competência de cada camada.

“Como assim, João? O que você quer dizer por ‘competência’?”

Simples: a cada camada, no micro de origem, compete incluir mecanismos (protocolos, informações de controle) na mensagem para que a mesma camada no destino consiga ler e interpretar a mensagem, não é mesmo? É a essa responsabilidade que as setas escuras se referem. Portanto, a camada de aplicação da origem quer falar com a aplicação do destino e

coloca, na mensagem, meios para isso; a camada de apresentação da origem quer falar com a camada de apresentação do destino e põe dispositivos para isso na mensagem. E assim por diante.

Vamos analisar a troca de um e-mail, finalmente, sob a visão do modelo OSI:

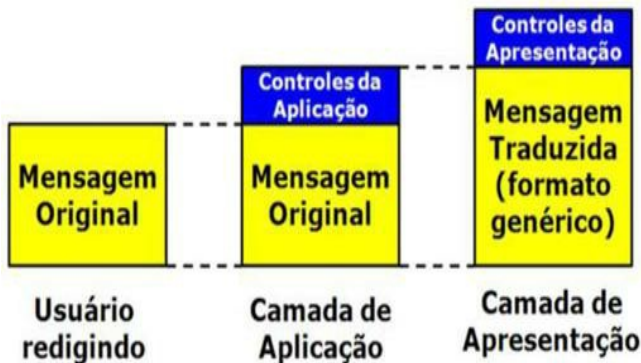
1. O usuário remetente do e-mail simplesmente o digita, em seu programa de e-mail preferido e clica no botão Enviar, para confirmar o envio – até aqui, dizemos que a responsabilidade é da camada de aplicação, com seus protocolos e serviços. Os protocolos da camada de aplicação simplesmente colocam “indicadores” (como “envelopes”) que permitirão à camada de aplicação do destino entender a mensagem como um todo.

Esse “envelope” é colocado, justamente, para que a camada de aplicação (o programa de e-mail) do usuário destinatário tenha condições de lê-la e retirá-la, a fim de ficar com a mensagem resultante e apresentá-la ao usuário. Também é importante dizer que esse “envelope” é apenas um conjunto de informações binárias (bits) que são adicionados à mensagem em si.

E a mensagem “envelopada” pela camada de aplicação desce...

2. Descendo, a mensagem é recebida pelos componentes da camada de apresentação, onde é devidamente modificada (seu conteúdo é reescrito de maneira diferente da original), por exemplo, elas são criptografadas, ou compactadas, ou escritas com outro conjunto de caracteres – ASCII, UNICODE etc.

Esse processo, como se sabe, tem por intuito traduzir a mensagem vinda da camada de aplicação (que pode estar em qualquer linguagem) para uma linguagem padrão genérica que pode ser entendida e transferida pelas redes. Novamente, isso é apenas utopia. Essa camada não existe na prática nos dias de hoje. Podemos ver esse “funcionamento”, até o ponto presente, na figura a seguir.



*Figura 8.61 – O funcionamento do envio da mensagem de e-mail (duas primeiras camadas).*

E a mensagem traduzida pela camada de apresentação desce...

3. ... E chega à camada de sessão. Aqui, são colocados controles (informações binárias) que indicam que aquela mensagem é uma (ou seja, inteira, maciça, única) e que, por causa disso, deverá ser transmitida em apenas uma sessão (um momento de comunicação).

Logo, nesta camada são inseridas informações que determinam e descrevem a sessão. Identificando toda mensagem como sendo parte daquela sessão. Doravante, os trechos de informação que forem divididos (pacotes, segmentos) vão carregar a marca daquela sessão.

E a mensagem daquela sessão... “Já sei, João: ‘desce...’” – é exatamente isso, caro leitor!

4. Descendo para a camada de transporte, aquela mensagem associada àquela sessão simplesmente passa por um processo traumático: ela é dividida! É na camada de transporte que a mensagem, não importando o seu tamanho original, se transforma em pequenos pedaços separados e autônomos.

“Os pacotes, João?”

Ainda não, caro leitor! Na camada de transporte, eles são normalmente conhecidos como **segmentos**. Só serão pacotes quando receberem os endereços de origem/destino válidos perante a estrutura da Internet (os endereços IP).

Na camada de transporte, os segmentos não sabem para onde vão, nem sabem como chegar lá. Mas sabem que farão parte de uma mesma mensagem, pois em cada segmento haverá uma indicação do tipo “1 de 10”, “2 de 10”, como já foi dito.

A figura a seguir mostra um resumo do que acontece enquanto uma mensagem desce da camada de apresentação até a camada de transporte (em que a mensagem é dividida em segmentos).



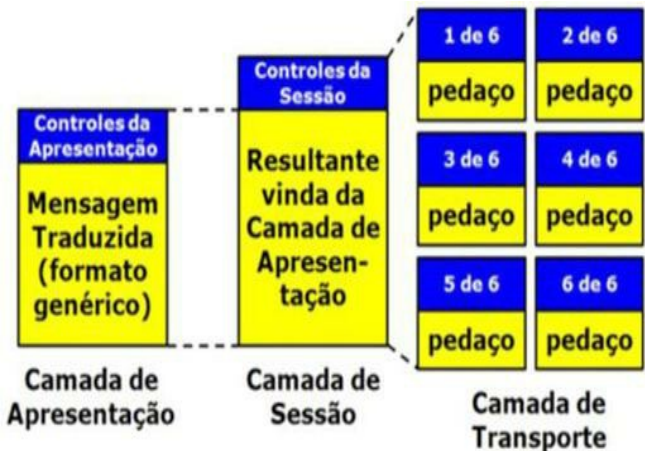


Figura 8.62 – A mesma mensagem: da camada de apresentação para a de transporte.

5. Descendo mais uma vez, os segmentos (vindos da camada 4) são transformados em pacotes. (Para isso, basicamente são adicionadas informações especiais que permitirão que os pacotes trafeguem pela estrutura da Internet.)

“Você quer dizer: são adicionados os endereços IP de origem e destino?”

Sim! Precisamente, caro leitor! As informações necessárias para que um pacote faça seu trabalho (ou seja, atravessar a estrutura das redes em direção ao destino) incluem os endereços IP de origem e destino.

“Então, João, a principal diferença entre ‘pacotes’ e ‘segmentos’, apesar de ambos serem trechos de informação, é que os pacotes possuem endereços?”

Sim! Isso mesmo... Veja a explicação um pouco mais detalhada:

5a. Os segmentos (trechos de informação da camada 4) sabem que fazem parte de algo maior (há controles que informam que eles são apenas “gotas d’água descobrindo que são mar azul...”), ou seja, são “peças de um quebra-cabeças” e têm consciência disso. Embora consigam “se juntar” na camada de transporte no destino, eles não conseguem trafegar até lá (porque não possuem nenhum tipo de endereçamento);

5b. Pacotes (pedaços da camada 3) não se entendem como parte de coisa alguma – são totalmente “autossuficientes” e “egocêntricos”. Os pacotes simplesmente sabem que têm (cada

um na sua) de atravessar uma estrutura de redes enorme até chegar ao ponto de destino (e têm condições de fazê-lo). Os pacotes possuem informações de endereçamento para permitir que saiam da origem e cheguem ao destino com perfeição.

Note, caro leitor, que neste ponto não há nem “sombra” da mensagem original. Todos os componentes dessas camadas apenas enxergam “trechos” que só serão unificados na camada de transporte no micro de destino. Nenhuma camada abaixo da camada de transporte faz ideia de que aqueles pedaços são partes de uma mensagem de e-mail.

Então, para finalizar o trabalho na camada 3, os outrora segmentos são inseridos em um “envelope” que contém endereços IP de origem e destino, conferindo-lhes o título de pacotes e dando-lhes condições de serem transferidos em uma estrutura composta de várias redes (como a Internet).

Mas, antes de sair dessa rede, esse pacote tem de ser capaz de atravessar a estrutura de enlace da rede de que o computador de origem faz parte. Ou seja, o pacote pode até ser capaz de atravessar várias redes (parabéns para ele!), mas precisa de uma ajudinha especial para trafegar pela rede da qual faz parte e chegar até o roteador (onde será, então, interpretado novamente como um pacote).

Então, novamente, a mensagem desce...

6. ... Para a camada de enlace (camada 2), onde é transformada em um quadro.

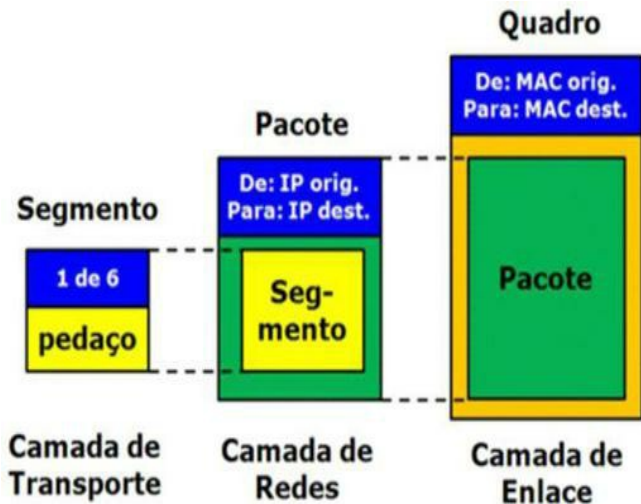
“OK, João. Lembro disso. Quadros são trechos de informação que só trafegam dentro de uma única estrutura de redes, pois utilizam como origem e destino os endereços MAC, não é?”

Eu não poderia explicar melhor, caro leitor. Parabéns!

Um pacote (que já é um “envelope” para um segmento) é, então, colocado em um “envelope” que chamaremos de quadro. Uma vez dentro do quadro, os endereços IP do pacote são simplesmente desconsiderados, porque só vão valer os endereços do novo envelope (os endereços MAC).

Isso é meio óbvio, não é? Pois se você colocar uma carta dentro de um envelope pequeno, escrever um endereço “A” nesse envelope e colocá-lo dentro de um envelope maior, dá para deduzir que o endereço “A” ficará inacessível (pois será conteúdo do envelope maior, e não endereço visível). Qualquer endereço “B” que você escrever no lado de fora do envelope grande será utilizado para fazer aquela correspondência trafegar, não é mesmo?

Então, um quadro “envolve” o pacote e lhe dá endereços que o tornam capaz de atravessar a rede local. A transição das camadas de transporte até enlace é vista na figura a seguir.



*Figura 8.63 – A mesma mensagem: da camada de transporte para a de enlace.*

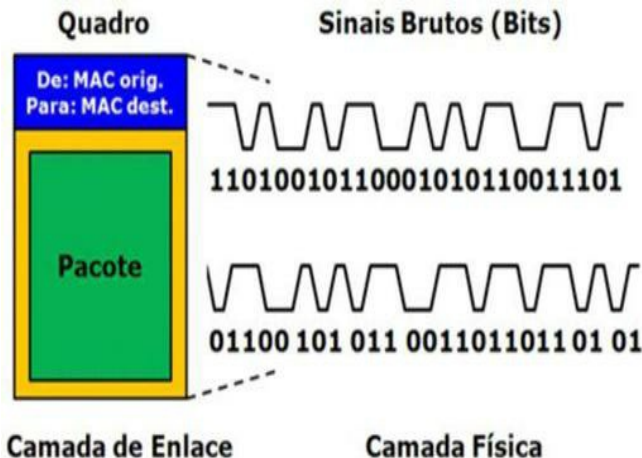
Agora chegou a “hora da verdade”. Os quadros precisam ser enviados, pela estrutura física da rede em questão (fios, hubs, conectores, antenas) e, para isso, precisam ser transformados em algo que possa ser transmitido.

Esse “algo” nada mais é que pulsos (elétricos, luminosos ou eletromagnéticos) que têm valores físicos determinados (uma frequência definida, um valor de tensão, uma corrente específica). Tudo isso depende da estrutura física da rede (a camada física) é por isso que dizemos que aquele quadro desce...

7. ... E se transforma em sinais brutos que podem ser transmitidos pela estrutura física da rede. É justamente assim que os equipamentos dessa camada (a camada física, ou camada 1) os enxergam. Em suma, todos os equipamentos da camada 1 apenas conseguem “ler” a informação (no nosso caso, o e-mail) como sendo sinais, meramente sinais brutos. Pulsos que representam 0 ou 1.

Mas, aí é que está... As informações (o e-mail) desde o momento em que foi escrito pelo usuário, lá “na camada de aplicação”, já eram um conjunto de zeros e uns (ou seja, bits).

Durante todo o trajeto, a mensagem sempre foi um conjunto de bits. Apenas foi vista (“enxergada”) sob outras ópticas a cada camada que descia.



*Figura 8.64 – Finalmente, a mensagem sendo vista na camada física como realmente é.*

Os sinais brutos que levam aquele quadro são, então, enviados pela estrutura da rede (fios, por exemplo) para chegar ao micro de destino, onde sobe, camada por camada, todas as sete:

1. Vamos supor que os sinais chegam pela camada física (fios) aos conectores da placa de rede. Por enquanto, eles não são entendidos como pacotes ou quadros, mas simplesmente como um punhado de sinais brutos (bits) vindos na forma de pulsos elétricos pelos fios.

Esses sinais são reunidos em grande quantidade...

2. ... E alçados à camada superior (camada de enlace – camada 2), onde aqueles bits são lidos como um quadro que contém em seu cabeçalho endereços MAC de origem e destino.

Os componentes da camada 2 (placa de rede, por exemplo) leem o endereço MAC contido no quadro para saber se lhe pertence (ou seja, para saber se o legítimo destinatário daquele quadro é realmente aquela placa).

Ao confirmar a posse do quadro, o cabeçalho do quadro é retirado, pondo à mostra o pacote, que é enviado à camada seguinte...

3. ... Que é justamente a camada de rede (camada 3). Nesta camada, o endereço IP do cabeçalho é lido. O micro, então, analisa se aquele endereço IP de destino é realmente dele (ou seja, o micro quer saber se aquele pacote é mesmo para si).

Em caso afirmativo, o pacote será processado por aquele micro, ação que começa com a retirada do cabeçalho do pacote, liberando o segmento...

4. ... Que será jogado para a camada 4 (camada de transporte) e minuciosamente analisado. (Para saber se é o único, se está com problemas, se será necessário pedir para que o micro de origem o reenvie, se está na ordem etc.)

Estando tudo certo com esse segmento, a camada de transporte passa a esperar pelos demais segmentos, para unificá-los e entregar o resultado desta unificação à...

5. ... Camada de sessão, que analisa o cabeçalho que a camada de sessão original havia posto. Com isso, a camada de sessão tem condições de saber se a sessão (transação) permanecerá aberta ou não para as próximas mensagens. Depois de “retirar” o cabeçalho posto pela camada de sessão original, a mensagem resultante é enviada...

6. ... À camada de apresentação (camada 6), que analisará o cabeçalho que a camada de apresentação (lá no micro de origem) colocou. Esse cabeçalho contém informações de como foi feita a tradução (de que linguagem para que linguagem) e, com isso, possibilitar a tradução no sentido inverso (de mensagem “genérica” para a mensagem “na linguagem oficial” da aplicação).

Por fim, a camada de apresentação, depois de traduzir a mensagem de volta ao seu formato original, a envia à...

7. ... Camada de aplicação (camada 7), que recebe a mensagem já em sua linguagem original, mas ainda com o cabeçalho inserido na camada de aplicação lá no micro de origem. Então, a camada de aplicação (ou seja, o software aplicativo que o usuário utiliza) vai “retirar” esse cabeçalho, deixando que o usuário leia apenas a mensagem. A mesma mensagem de e-mail que fora redigida por seu remetente.

Ufa! E aí, leitor? Tudo tranquilo?

“Mais ou menos, João. Ainda estou digerindo isso tudo.”

É só ler novamente, espero! Garanto a você uma coisa: levei um ano para entender completamente os modelos de camadas, porque, simplesmente, não entendia para que servia. (É um “bloqueio” meu, um “tabu”... Acho que é “tapadice” mesmo... Eu só consigo entender algo se souber para que ele funciona). Meus agradecimentos sinceros aos professores **Juliana Diniz** e **Obionor Nóbrega**, meus mestres em redes na faculdade. Eles esclareceram tudo e permitiram que esta explicação fosse possível.

Mas ainda tem mais: o modelo OSI não é o mais usado (nem o mais cobrado nas provas) atualmente. Esse modelo tem algo de teórico, utópico, pouco usual. Isso se deve ao fato de a Internet se basear em um modelo anterior ao OSI – o modelo de camadas TCP/IP, que conheceremos agora.

### 8.9.2. Modelo de camadas TCP/IP

TCP/IP é o nome dado a um conjunto de protocolos (ou “pilha” de protocolos). Sua importância é incontestável. A Internet baseia sua comunicação nessa pilha de protocolos. Ou

seja, todos os computadores da Internet (hoje, cerca de 1 bilhão) “falam” os protocolos contidos na *pilha TCP/IP*.

É fácil entender também que, para se tornar padrão, o funcionamento da Internet (incluindo seu conjunto de protocolos) precisou ser padronizado, esquematizado, normatizado.

“Ou seja, João, foi necessário escrever um modelo de camadas para ele, não é?”

Sim! Exatamente! O modelo de camadas *TCP/IP*. Que, inclusive, foi proposto e aprovado antes do OSI. (O OSI foi uma tentativa de “unificar” todos os modelos de camadas até então existentes.)

O nome TCP/IP é formado pelo nome dos dois mais importantes protocolos deste conjunto: o *TCP* (Transmission Control Protocol – Protocolo de Controle da Transmissão – pertencente à camada de transporte) e o *IP* (Internet Protocol – Protocolo de Inter-redes – localizado na camada de rede).

O modelo de camadas TCP/IP é formado por quatro ou cinco camadas.

“Quatro ou cinco? Como assim, João?”

É o seguinte, caro leitor. Como não há um padrão documental sobre isso (assim como no OSI), visto que o modelo TCP/IP é, meramente, um “acordo”, uma “política de boa vizinhança”, alguns autores (conceituados, inclusive, como Douglas E. Comer) o desenham com cinco camadas, e outros autores (como Andrew Tanenbaum, o “papa” em redes) preferem desenhá-lo com apenas quatro camadas.

“E nas provas, João?”

Não sei! Sinceramente, pode aparecer qualquer um deles! O que posso afirmar, com certeza, é que será exigido o conhecimento nos protocolos que o compõem e não exatamente na quantidade de camadas que o formam.

Uma coisa é certa: apesar de ser semelhante ao ISO/OSI, o modelo TCP/IP não é derivado deste e, portanto, camadas homônimas nos dois modelos podem, sim, apresentar objetivos e características diferentes entre si, o que torna o estudo do modelo TCP/IP relativamente desligado do estudo do OSI.

Eis os modelos TCP/IP de cinco e quatro camadas:



Figura 8.65 – Modelos de camadas TCP/IP.

Note três características semelhantes nos dois modelos em relação ao modelo OSI:

- As camadas de **apresentação** e **sessão** sumiram! As funções que, no modelo OSI, são responsabilidade dessas duas camadas foram assimiladas pela camada de **aplicação**. Portanto, lembre-se de que nas comunicações da Internet, o estabelecimento de sessões e a tradução da mensagem (como criptografia e compactação) são responsabilidade da camada de aplicação.
- A camada de redes (camada 3 no OSI) passou a se chamar **Camada de Inter-Redes**. Isso

é bom porque explicita o objetivo dessa camada: a ligação entre redes distintas.

c. A camada de enlace (camada 2 no OSI) passou a ser chamada de ***Camada de Interface de Redes***.

A principal diferença entre os modelos é que os defensores de quatro camadas apenas “interpretam” que as camadas 1 e 2 são uma só. Ou seja, esses autores definem que não há a camada física e a camada de interface de redes, mas apenas uma que acumula a função das duas.

Essa “possibilidade” de interpretação em duas formas tão distintas se deve ao fato de, na verdade, o modelo TCP/IP só definir a existência e o funcionamento de componentes nas três camadas superiores.

“Peraí, João! Tá querendo me dizer que o modelo TCP/IP só apresenta componentes nas camadas de aplicação, transporte e inter-redes?”

Sim! Perfeitamente, caro leitor! Afinal, TCP/IP é um conjunto de protocolos (e protocolos são programas). Em um modelo de camadas que se baseia na estrutura de um conjunto de protocolos, ou seja, em um conjunto de programas, a definição ou exigência quanto a componentes físicos (camadas física e de interface de rede) não seriam muito adequadas.

Então, fique ciente disto: o modelo TCP/IP só estabelece padrões e definições nas três camadas superiores. Isso quer dizer que o modelo de camadas TCP/IP “não se importa” com o que existe nas camadas física (1) e de interface de rede (2).

Com isso, chegamos a uma característica forte e importante na Internet: não importa quais são as estruturas físicas de rede que ligam os computadores em uma rede. Se essa rede possuir os mesmos protocolos das camadas superiores (inter-redes, transporte e aplicação), ela poderá se ligar à Internet.

#### **8.9.2.1. O “copo de liquidificador universal”**

Uma das inúmeras comparações que faço em sala de aula é a de que a pilha de protocolos TCP/IP é como um “copo de liquidificador universal”. Acompanhe a história:

Imagine que um vendedor (daqueles de porta em porta) chega à casa de uma senhora e diz:

“Boa tarde, minha senhora, estaria interessada em um copo de liquidificador universal? Ele simplesmente encaixa em qualquer base, não importa a marca, idade ou nacionalidade do seu liquidificador.”

“Não, obrigada” – responde a cliente –, “meu liquidificador está com o copo quebrado há mais de cinco anos, mas, como ele é importado da Ucrânia, nunca consegui achar um copo no formato adequado.”

“A senhora não me entendeu!” – insiste o vendedor – “O copo que vendo é universal! Ele se adaptará a qualquer base. Qualquer base mesmo! Na verdade, não importa se a senhora vai me mostrar a base ou não. Eu sei que o copo vai servir!”

E assim é o TCP/IP. É o “copo universal”! Ele serve para permitir a comunicação entre diversas redes diferentes, não importa qual é a tecnologia (Token Ring, Ethernet, Wi-Fi, WiMAX etc.) e, consequentemente, a estrutura física das redes interligadas. Não importa mesmo! O TCP/IP vai funcionar!

“Mas, João, por que ‘não importa’ se as redes que vão se interligar têm ou não têm a mesma



estrutura? Pensei que ter a mesma arquitetura e estrutura física fosse uma condição necessária para que as redes se comunicassem bem.”

Caro leitor, é só refletir um pouco: em uma rede única (vários micros ligados por hubs ou switches), o que importa é o protocolo da camada 2 (enlace, ou interface de redes). Portanto, todos os equipamentos (placas de rede, switches, pontes) têm de concordar quanto à arquitetura e ao protocolo de acesso ao meio. Ou seja, todos têm de ser feitos na mesma arquitetura de rede.

Em ligação entre redes distintas, outro personagem aparece na jogada: o roteador. Sim! O roteador, que é um equipamento da camada 3, não se importa se as duas redes são da mesma arquitetura ou de arquiteturas diferentes, pois a comunicação entre elas se dá por meio de pacotes (pedaços de informação da camada 3) e, por isso, até mesmo a regra de endereçamento se dá alheia à tecnologia empregada nas redes em questão.

Ou seja, em uma comunicação entre redes (inter-networking, ou simplesmente inter-net, termo que originou a Internet), não importam as tecnologias, mas somente o que está da camada 3 (inter-redes) para cima!

### 8.9.2.2. As camadas do modelo TCP/IP

Eis um resumo das camadas (e, claro, dos componentes delas) no modelo de camadas TCP/IP:

- **Camada 5 – Aplicação:** nesta camada estão os protocolos de mais alto nível, aqueles que realizam tarefas diretamente em contato com os usuários: FTP, SMTP, HTTP, POP, IMAP, DNS, TELNET, NNTP etc.

Esses protocolos estão intimamente ligados às diversas tarefas (serviços) que podemos utilizar na Internet. (Normalmente, cada protocolo está associado a um serviço diferente.)

- **Camada 4 – Transporte:** estão localizados, nesta camada, os protocolos responsáveis pela comunicação fim a fim entre as máquinas envolvidas. Os protocolos da camada de transporte são: TCP e UDP.

Os protocolos da camada de aplicação precisam dos protocolos da camada de transporte. Algumas aplicações (programas) usam o UDP, mas a grande maioria dos protocolos localizados na camada 5 usa o TCP como protocolo de transporte.

- **Camada 3 – Rede (ou Inter-Redes):** apresenta protocolos que realizam processos de roteamento e tradução de endereços para que a conexão entre os dois computadores seja efetuada. Fazem parte desta camada os protocolos IP, ICMP, IGMP, ARP e RARP.

Desses vários protocolos, o mais importante (considerando todas as camadas da pilha) é, sem dúvidas, o IP. Todos os protocolos das camadas superiores precisam do IP, que é o responsável direto pelo endereçamento dos micros e pelo roteamento dos pacotes através da estrutura das redes.

*Sem IP, não há comunicação.*

- **Camada 2 – Enlace (ou Interface de Rede):** o modelo TCP/IP não se “mete” com ela, porque não se “importa” com o tipo da arquitetura da rede (ou seja, para o TCP/IP, não há necessidade de saber se a rede é Ethernet ou Token Ring.). O termo mais “polido” para esse caso é “o modelo TCP/IP não especifica padrões de equipamentos nem protocolos para a camada de enlace”.

Lembre-se do “copo de liquidificador universal”!

Mas, para essa camada, continuam valendo as funções atribuídas a ela pelo modelo OSI.

- **Camada 1 – Física:** o modelo TCP/IP também não especifica padrões para a camada física. (Me diz se esse texto não ficou bonito.). Em outras palavras: a rede pode ser montada com qualquer tipo de cabo, fio, fibra etc., o TCP/IP não se “importa” com isso.

As características e funções da camada física são descritas no modelo OSI.

## 8.10. Protocolos da pilha TCP/IP

Vamos começar uma análise a respeito dos principais protocolos que formam o conjunto de protocolos TCP/IP. Esses protocolos são apenas “linguagens” que permitem a comunicação entre computadores. Como eles têm funções definidas por um padrão de camadas (visto anteriormente), vamos estudá-los por essa classificação (separados por camadas).

## 8.11. Protocolos de rede

Os protocolos de rede são os de mais “baixo nível” do conjunto TCP/IP. A função genérica dos protocolos dessa camada é criar meios de a mensagem trafegar pela estrutura das inter-redes, facilitando a decisão do melhor caminho a ser tomado pela mensagem para a chegada ao destino. Os protocolos descritos como pertencentes à camada de rede são:

- IP;
- ICMP;
- ARP;
- RARP.

### 8.11.1. Protocolo IP

O protocolo IP (Internet Protocol – Protocolo de Inter-Redes) é o mais importante da pilha TCP/IP, tanto que, junto com o TCP, dá nome a ela. As duas funções do protocolo IP são endereçar as estações de origem e destino (ou seja, dar a cada um deles um endereço) e rotear as mensagens entre elas (rotear é “definir a rota”).

Como é um protocolo da camada 3 (inter-redes), o IP é responsável por manipular pequenas unidades de informação chamadas *pacotes*. Tais pacotes também podem ser chamados de datagramas IP, ou simplesmente *datagramas*.

Um pacote IP é tipicamente formado por um conjunto de bits com tamanho máximo especificado. Nesse pacote estão duas áreas distintas (aliás, como em qualquer unidade de transporte de dados, seja um segmento ou um quadro): a área de cabeçalho e a área de dados do pacote (também chamada de payload).

Na área de dados está encapsulado (“dentro de um envelope”) um segmento vindo da camada superior (camada de transporte) – já vimos isso na explicação das camadas, não foi?

No cabeçalho do pacote IP (aquela parte do pacote que contém as informações de controle daquela camada) estão informações como: endereço IP de origem, endereço IP do destino, tempo de vida (TTL – time-to-live) do pacote, protocolo superior, entre outras. Dá uma olhada num pacote “desenhado ludicamente” e sua “verdadeira forma”.

# Pacote IP



```
01010001010101110
10101010010101111
11010101110101010
10101010010101010
10101010010101010
10101000000000101
01010101010101010
10101010101011111
11110101101011010
01010111010111000
```

Como “representamos”

Como ele é realmente

Figura 8.66 – Um pacote IP – lembre-se: é sempre um “punhado” de bits, apenas...

Vale salientar, também, que o pacote IP não é exatamente como mostrado na figura anterior. Há outros campos no cabeçalho, como tamanho do pacote, checksum do cabeçalho (para detecção de erros), tipo de serviço, entre outros.

“Por que não mostrar todos os campos, então, João?”

Porque acredito que para o entendimento do funcionamento do IP, eles não são necessários. Além disso, são muitos campos e muito detalhados. Seu estudo é necessário apenas para o pessoal da área de Informática (que não é o objetivo deste livro, só lembrando, embora muitos o utilizem!).

Cada pacote IP contém, em seu cabeçalho, entre outras, as seguintes informações:

- **Endereço IP de Destino:** é o endereço (formado por 32 bits – ou seja, 32 “zeros” e “uns”) que determina a máquina (estação) destinatária daquele pacote.
- **Endereço IP de Origem:** é o endereço (de mesmo tamanho, ou seja, 32 bits) que indica qual é a máquina remetente daquele pacote.
- **TTL (Time-to-Live – Tempo de Vida):** é um número de segundos (ou saltos – hops) que o pacote deve “viver” para atravessar a Internet. A cada passagem por um roteador (hop, ou salto), esse número é diminuído. Caso ele atinja 0 (zero), o próximo roteador simplesmente “descartará o pacote”, em vez de retransmiti-lo. TTL é, em suma, uma espécie de “data de

validade” do pacote.

- **Protocolo:** é um número (de 8 bits de tamanho) que identifica o protocolo encapsulado no pacote (ou seja, identifica o “conteúdo do envelope”). Esse campo serve para que o computador/roteador consiga saber se aquele pacote será enviado à camada de transporte (se for TCP ou UDP) ou se o próprio roteador poderá manipulá-lo (por exemplo, quando o protocolo encapsulado for ICMP ou IGMP, ambos da camada 3).

- **Checksum do Cabeçalho:** é um número (de 16 bits) que atua como “um resumo” matemático do cabeçalho. A cada hop (salto – passagem por um roteador), o cabeçalho vai mudar em algum ponto (nem que seja somente no TTL). Depois de mudado o cabeçalho, é recalculado outro checksum (em português, soma de verificação) que resumirá matematicamente aquele cabeçalho. Se você já estudou segurança, é como um “hash” do cabeçalho. Serve para que os roteadores identifiquem se há erros na transmissão daquele cabeçalho (detecção de erros).

- **Comprimento (Length):** esse campo identifica o tamanho que o pacote tem, em bytes. Um pacote IP pode ter 576 bytes (no mínimo) e 65.536 bytes (64 Kilobytes) no máximo.

“João, preciso saber esses dados? Esses ‘campos’ no cabeçalho IP?”

Só falta uma coisa para ser lembrada: o protocolo IP é considerado não-orientado a conexão. Isso significa que o protocolo IP não se preocupa em estabelecer conexões prévias entre origem e destino para poder transmitir. Nem se preocupa se o pacote chegou ou não. Nem exige qualquer tipo de confirmação do destinatário.

“Peraí, João! Agora é demais. Como ele não exige confirmação? E como ele vai saber que o pacote chegou?”

A resposta é simples, caro leitor: ele não precisa saber! O protocolo IP não “tá nem aí” para o pacote! Se ele chegar, bom... Se não chegar, bom também!

“E por que esse descaso?”

Não é descaso. É apenas desnecessário, visto que a camada de transporte (que é responsável por incluir controles que darão por falta de um segmento) é que deve fazer isso. Portanto, o trabalho de se “estressar” pela falta do pacote não é do IP! É do TCP. Portanto, se essa função já é realizada na camada superior, não é necessário fazê-la na camada 3.

Não se preocupe, que veremos TCP mais adiante. Agora é o momento de entender tudo sobre endereço IP. Aproveite...

### 8.11.2. Endereço IP

Endereço IP é o endereço numérico que identifica qualquer conexão feita a uma estrutura de inter-redes baseada em TCP/IP. Ou seja, endereço IP é o endereço usado na camada 3 (inter-redes) do modelo de camadas TCP/IP.

“João, o que você quis dizer como ‘identifica qualquer conexão’ não seria ‘identifica qualquer máquina (ou estação)’? E por que o endereço IP é tão importante que mereça ser estudado a fundo?”

Em primeiro lugar, caro leitor, o IP não identifica uma máquina. Se um computador, por exemplo, possuir duas placas de rede ligadas simultaneamente a uma mesma rede, cada uma delas possuirá um endereço IP associado. Portanto, a máquina em si teria dois endereços IP.

Em segundo lugar, sobre a importância do endereço IP. Como a Internet que conhecemos é baseada no modelo de camadas TCP/IP, e, consequentemente, em seus protocolos, então o endereço IP é a forma oficial de endereçamento na Internet.

O endereço IP é um número binário (aliás, como tudo na comunicação digital) formado por 32 bits. Em suma, um endereço IP é exatamente assim:

11001000111110010000110111101100

A pergunta que não quer calar é: quem se lembraria de um desses? Eu imagino um administrador de redes dizendo: “Eita, deu problema no computador da secretária do financeiro... qual é mesmo o endereço de lá? 01...11...10... não, não, é 01111110... Ahhh!”

Por ser realmente muito complicada sua memorização (e utilização), os endereços IP não são representados no seu formato puro. Usa-se uma forma de notação em que se divide o endereço IP em **4 grupos de 8 bits** (1 byte cada, ou, como costumamos chamar, 1 *octeto*.)

11001000.11111001.00001101.11101100

(Esses pontos não existem nos endereços IP de verdade. São simplesmente para demonstrar a separação.)

Depois de separarmos os grupos de octetos, **convertemos** esses octetos para **números decimais**, resultando em algo assim:

200.249.13.108

Essa forma de “representação” é chamada notação decimal separada por pontos.

“Calma lá, João... Como 11001000 foi se transformar em 200?”

Através de um processo simples de conversão de binário (zeros e uns) para decimal (base numérica que usamos em nossa matemática). Esse processo está descrito com perfeição no capítulo sobre **Aritmética Computacional**, no final deste livro. Se quiser aprender como fazer, pare a leitura deste capítulo exatamente agora e vá ler o capítulo que citei. Depois volte aqui e passará a ver, tenho certeza, esses processos com outros olhos.

Voltando ao assunto, cada octeto é representado por um número decimal, que poderá variar entre 0 (que em binário seria 00000000) e 255 (que é 11111111). Então, podemos dizer por dedução, que o endereço IP é um endereço numérico binário representado de forma decimal por quatro números, separados por pontos, que podem, cada um, assumir qualquer valor entre 0 e 255.

Um computador que vai se ligar à Internet, ou mesmo apenas a uma rede local que usa TCP/IP como pilha de protocolos, precisa ter endereço IP. Se um computador não possuir endereço IP, não poderá enviar nem receber pacotes. Estará, portanto, ilhado. Não conseguirá se conectar à rede.

“João, basta que um computador saiba seu próprio endereço IP para estar apto a trocar dados na rede? Ou tem mais alguma coisa?”

Já que você perguntou, caro leitor...

### 8.11.3. Parâmetros IP

Para que um computador ligado a uma rede que usa TCP/IP seja capaz de se conectar a uma rede a fim de trocar informações com outros computadores, é necessário que ele conheça duas informações básicas:

a. Seu próprio endereço IP;

b. A máscara de sub-rede da rede da qual ele faz parte.

E tem mais! Essas duas informações permitem que o micro se ligue a outros *em uma só rede*. Se você quiser que o micro se ligue na Internet (ou seja, com várias redes distintas), ele deverá conhecer uma terceira informação:

c. O endereço IP do gateway padrão (ou seja, do roteador) da sua rede.

Essas informações são genericamente conhecidas como parâmetros IP e são necessárias para que qualquer computador se ligue à Internet.

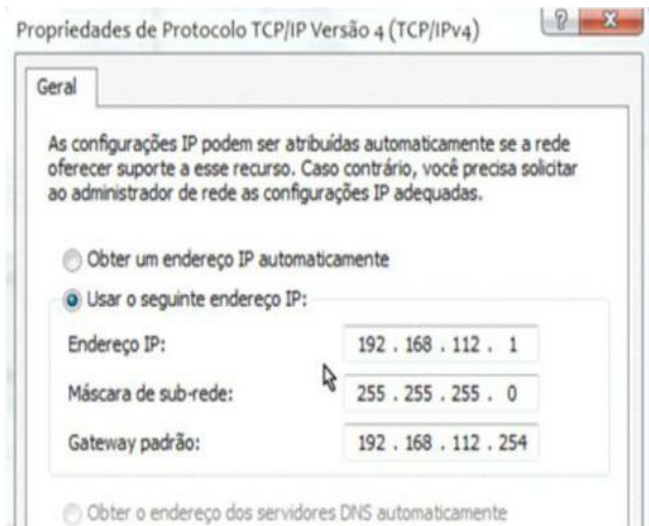


Figura 8.67 – Configurando os três parâmetros IP em um micro com Windows Vista.

#### 8.11.3.1. Endereço IP do próprio micro

Quanto ao endereço IP do próprio computador, não há o que discutir, não é? Quero dizer: se você não soubesse qual é o seu nome, quando alguém gritasse por você, não atenderia porque não identificaria o chamado, não é?

Então é isso... Se alguém grita, na rua: “EI, JOÃO!!!” eu vou olhar, simplesmente porque sei que meu nome é João. Se eu não soubesse, nem olharia! Isso me permite concluir que um computador precisa saber seu próprio endereço IP.

E quanto ao formato do endereço IP, não há problemas, não é?

Mas o que são os outros dois parâmetros?

### 8.11.3.2. Endereço IP do gateway padrão

É apenas o endereço IP do roteador daquela rede. Todo computador precisa saber qual é o endereço do roteador que o serve. Isso é necessário porque quando um computador perceber que vai transmitir um pacote para outra rede (não para a rede da qual ele faz parte), ele enviará o pacote àquele que poderá enviá-lo a outras redes: *o roteador*.

Portanto, para que um micro consiga se comunicar na Internet, ele tem de saber o endereço IP do seu roteador (gateway padrão ou “portão padrão”). Caso um micro não saiba essa informação, mas saiba seu próprio IP e a máscara de sub-rede, ele conseguirá se comunicar internamente (com outros micros na mesma rede), mas não na Internet.

### 8.11.4. Máscara de sub-rede

A máscara de sub-rede também é, a exemplo do endereço IP, uma informação binária de 32 bits (32 “zeros” e “uns”). A máscara de sub-rede também pode ser representada como um conjunto de quatro octetos decimais separados por pontos, como o próprio endereço IP.

Porém, existe uma coisa muito peculiar na máscara de sub-rede: ela é formada por 32 bits, sendo que *inicia* com um bloco ininterrupto de *1 (uns)* seguido de um bloco ininterrupto de *0 (zeros)*. Sem alternância.

Ou seja, isto aqui é uma máscara:

111111111111111111111111000000000000

E isto aqui não é uma máscara (mas poderia ser um endereço IP de algum micro):

11001100111100010101011101011110

“Ei, João. Sendo assim, quando convertermos essa máscara para decimal, os octetos resultantes não poderão apresentar todos os valores entre 0 e 255, não é mesmo?”

Exatamente. Parabéns, leitor! A máscara de sub-rede, quando apresentada em sua forma pura (binária), é representada como uma sequência de uns seguida de uma sequência de zeros, como vimos, e isso limita o formato decimal da máscara para alguns valores.

Só podem ser octetos em uma máscara em decimal os números:

255 – porque é 11111111 em binário;

254 – porque é 11111110;

252 – porque é 11111100;

248 – porque é 11111000;

240 – porque é 11110000;

224 – porque é 11100000;

192 – porque é 11000000;

128 – porque é 10000000; e

0 – porque é 00000000;





Um endereço IP não serve apenas para identificar uma estação em si (ou uma conexão à Internet). Inerente ao endereço IP, existe uma informação que identifica a rede da qual aquela estação faz parte.

É que o endereço IP pode ser visto como um “nome completo” ou pelo menos daqueles nomes que se encontram em passagens de ônibus e avião: *Carvalho/João* ou *Silva/Eduardo*.

Então, o endereço 200.234.44.112 não serve para identificar somente um micro. Nesse endereço há a identificação de duas coisas: do micro em si (ID do host, ou ID da estação) e da rede (ID da rede). Resta saber qual é o ID da rede e qual é o ID do host dentro do endereço IP. (Atenção – é ID mesmo! ID vem de *Identificador*)

Que tal se perguntássemos assim: no endereço 200.234.44.112, quais octetos representam a rede e quais octetos representam o micro em si? Seria o mesmo que perguntar: no nome João Antonio César Carvalho, quais os nomes que representam a família e quais os nomes que representam o indivíduo? Difícil saber.

A máscara faz isso, caro leitor! A máscara atua como a / (barra) em Carvalho/João Antonio, permitindo que se possa determinar quem é família (Carvalho) e quem é indivíduo (João Antonio). Só que a máscara faz isso com endereços IP.

Vamos aplicar uma máscara em um endereço IP usando a notação de decimais separados por pontos. Para isso, porém, é bom que se saiba que só será possível fazer os cálculos com três máscaras apenas (aquelas que usam os octetos completamente preenchidos ou por 1, ou por 0). Seriam elas:

- 255.0.0.0 (máscara dos endereços Classe A);
- 255.255.0.0 (máscara dos endereços Classe B);
- 255.255.255.0 (máscara dos endereços Classe C);

“Ô, João, e esse negócio de Classe? O que é?”

Daqui a pouco veremos isso. Vamos continuar com o entendimento das máscaras sem essas informações.

Para todas as demais máscaras de sub-rede possíveis, o cálculo que vamos aprender agora só será possível se convertermos as máscaras e os endereços IP para binário. (Depois ensino isso... Prometo!)

#### 8.11.4.2. Analisando a máscara classe C

Então, vamos lá. Endereço IP 192.168.214.123 e máscara de sub-rede 255.255.255.0. O que posso fazer com esses dados? Analise-os verticalmente (um em cima do outro).

**192.168.214.123**

**255.255.255.0**

Aqueles octetos do endereço IP que coincidirem, em posição, com os octetos 255 da máscara são os que representam a rede. Por sua vez, os octetos do endereço IP que coincidirem com os octetos 0 da máscara representam o micro (o indivíduo).



Figura 8.68 – Analisando o IP 192.168.214.123 na máscara 255.255.255.0.

Então, de uma maneira bem “rústica” e “acústica”, o nosso computador mostrado na figura pode ser identificado como o **micro 123**, pertencente à rede cujo “prefixo” é **192.168.214**. Ou seja, em uma máscara classe C, os três primeiros octetos representam o ID da rede e apenas o último octeto representa o ID do micro. Simples, não?

E tem mais. Se outro micro qualquer possuir a mesma máscara e os mesmos três primeiros octetos, esse outro micro pertence à mesma rede que o micro do nosso exemplo. Vamos ver?

**192.168.214.123**

**192.168.214.30**

**192.168.214.249**

**255.255.255.0 (máscara de sub-rede)**

Todos esses micros acima fazem parte da mesma rede. E lembre-se de que todos os micros da mesma rede têm de ter a mesma máscara de sub-rede definida. Observe que os octetos do ID da rede são sempre os mesmos para todos os micros naquela rede (óbvio, né?), o que obriga que, de um micro para outro, só varie o último octeto.

Ao que eu pergunto: quantos micros são possíveis em uma rede qualquer cuja máscara de sub-rede é 255.255.255.0 (classe C), caro leitor?

“256 micros! Pois como só quem varia de um micro para o outro é apenas o último octeto, e ele pode variar de 0 (zero) a 255. São 256 combinações possíveis!”

Mais ou menos, leitor... Seu raciocínio está perfeito! Mas não contei um segredinho: dois endereços são proibidos – o primeiro e o último!

#### 8.11.4.3. Endereço IP da rede e endereço IP de broadcast

Quando a estrutura de endereçamento de uma rede (ou seja, sua máscara de sub-rede e seu prefixo) é definida, dois endereços nunca (nunca mesmo) poderão ser usados para identificar um micro.

O primeiro endereço possível de se construir (usando os dados do nosso exemplo, seria 192.168.214.0) não é usado para identificar micros porque é usado para identificar a rede em si.

É um endereço hipotético que não tem função para a comunicação na rede, mas que a representa.

Portanto, o micro 192.168.214.123 não pertence à rede 192.168.214. Dizemos que ele pertence à rede **192.168.214.0**! Logo, o *primeiro endereço em uma rede é o endereço da rede em si*.

O outro endereço que não pode ser usado para identificar micros na rede *é o último possível*, ou seja, 192.168.214.255, tomando como base o nosso exemplo. O último endereço é chamado *endereço de broadcast* e serve para enviar uma mensagem a todas as estações daquela rede (ou seja, a todas as estações que comecem seus IPs por 192.168.214).

Portanto, em uma rede classe C (esse termo “classe C” significa que a rede usa a máscara 255.255.255.0), podemos ter até 254 computadores conectados porque podemos dar até 254 endereços IP (256 combinações possíveis menos 2 proibidos).

“Então é só isso, João? Os três primeiros octetos de um endereço IP sempre representam a rede e o último sempre representará o micro?”

Não! Não é sempre! Isso aconteceu porque a máscara usada definiu assim! Se fosse outra máscara, o caso seria diferente. Vamos analisar?

#### 8.11.4.4. Analisando a máscara classe B

Uma máscara de sub-rede de classe B tem os dois primeiros octetos representando a rede e os dois últimos octetos representando o micro (ou seja, 255.255.0.0).

*203.140.3.129 (endereço IP do micro que analisaremos)*

*255.255.0.0 (máscara de sub-rede classe B)*

Podemos dizer que esse é o micro “**3.129**” (“três ponto cento e vinte e nove” e não “três mil cento e vinte e nove”, como você poderia ler) dentro da rede cujo prefixo é “203.140”.

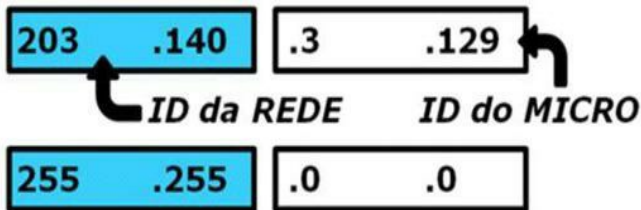


Figura 8.69 – Analisando o IP 203.140.3.129 na máscara 255.255.0.0 (classe B).

“Ei, João, posso chamar a rede de 203.140.0.0 – usando o primeiro endereço dela?”

Sim! Perfeito! É exatamente isso! O primeiro endereço é sempre aquele que representa a rede. Portanto, a rede cujo prefixo é 203.140 e cuja máscara é 255.255.0.0 é chamada de rede 203.140.0.0. (Logo se percebe que esse endereço não pode ser usado para micros, pois é o

primeiro.) Mas cadê o último?

“João, seria 203.140.255.255, porque os dois octetos finais variam de micro para micro?”

Sim! É exatamente isso! O último endereço (que vai servir como endereço de broadcast) de uma rede classe B tem os dois últimos octetos como sendo 255.

Note que, usando a máscara **255.255.0.0**, os endereços

**203.140.3.129**

**203.140.188.2**

**203.140.0.255**

**203.140.1.0**

**203.140.123.122**

pertencem à mesma rede (e são válidos para serem usados em micros, pois não são nem o primeiro nem o último endereços da rede).

Agora, caro leitor, a pergunta novamente: quantos micros são possíveis em uma rede com essa máscara de sub-rede?

“Ah, João... Como os dois primeiros octetos serão sempre os mesmos em todos os micros da rede, então somente os dois últimos octetos podem variar de micro para micro. Como cada octeto é independente um do outro e pode variar 256 vezes, isso vai dar 256 x 256 possibilidades de combinação, ou seja, 65.536 combinações. Ah! Claro... Sem o ‘0.0’ e o ‘255.255’, são 65.534 endereços para computadores possíveis.”

Perfeitamente! É exatamente isso! 65.534 computadores podem ser conectados a uma rede classe B. Vamos analisar uma rede classe A?

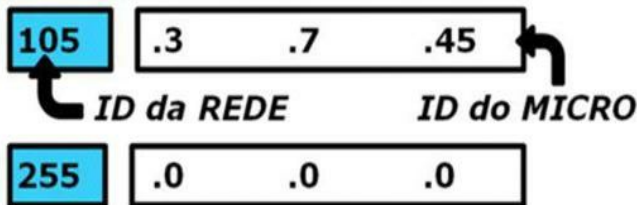
#### 8.11.4.5. Analisando a máscara classe A

A máscara de sub-rede classe A é aquela (dentre as três que vimos) que permite as maiores redes de computadores, pois apenas o primeiro octeto representa o ID da rede e os outros três octetos representam o ID do host (ou seja, 255.0.0.0).

**105.3.7.45 (Endereço IP do micro que estamos analisando)**

**255.0.0.0 (Máscara de sub-rede classe A)**

Sem dúvidas, podemos concluir que este seria o micro “3.7.45” dentro da rede “105”.



Não custa concluirmos, caro leitor, que o primeiro endereço (que será usado como “endereço da rede”) é 105.0.0.0 e que o último endereço (que será usado como endereço de broadcast) é 105.255.255.255.

Veja alguns computadores pertencentes à mesma rede classe A do nosso exemplo:

**105.3.7.45**

**105.2.234.255**

**105.23.0.0**

**105.214.249.254**

OK, OK... Agora a pergunta: quantos micros (hosts) pode haver em uma rede classe A, caro leitor?

“Já tá ficando repetitivo, João! Em uma rede classe A, apenas o primeiro octeto representa a rede; portanto, apenas ele ficará fixo (idêntico) em todos os micros da rede. Os três octetos finais podem variar. Como são três números que podem ir de 0 a 255, são 256 x 256 x 256 possibilidades. Ou seja, 16.777.216 combinações possíveis.”

Está se esquecendo de algo, leitor?

“Claro! 16.777.216 combinações menos os dois endereços proibidos (o primeiro – que é o endereço da rede – e o último – que é o do broadcast). Portanto, uma rede classe A pode ter até 16.777.214 micros.”

Exato. Para facilitar os seus cálculos, que tal entender uma fórmula simples? Pense que **K** é o número de *octetos 0 (zero) da máscara*, que pode ir de 1 (na classe C) até 3 (na classe A). Sendo **M** o número de *micros (hosts)* possíveis na rede, a equação é a seguinte:

$$M = 256^K - 2$$

Ficou claro isso, leitor? Agora tem mais... Vamos analisar a utilização prática (e comercial) das classes de endereços IP na Internet.

#### 8.11.5. Classes de endereços IP na Internet

A IANA (Internet Assigned Numbers Authority – Autoridade de Números Designados na Internet) e a ICANN (Internet Corporation for Assigned Names and Numbers – ou Corporação para Nomes e Números Designados na Internet) são os órgãos responsáveis por estabelecer os padrões de endereçamento usados na Internet, entre eles, a definição dos endereços IP e suas classes.

Segundo as regras criadas pela IANA, existem cinco classes de endereços IP (não somente A, B e C) e as regras mostradas aqui são seguidas exclusivamente para a atribuição de endereços IP em micros que vão se ligar diretamente à Internet (ou seja, em endereços IP que estão visíveis perante a Internet). Vamos a elas:

As redes classe A foram todas vendidas para grandes provedores de acesso e empresas de telecomunicações. São, ao todo, 126 redes classe A possíveis na Internet. Cada uma dessas redes pode ter, no máximo, 16.777.214 micros (hosts) – isso, claro, por causa da máscara das redes classe A (255.0.0.0).

“João, a parte dos micros eu entendi. Mas por que só pode haver 126 redes classe A na

Internet? E por que todas já foram vendidas?”

Sim... Todas as redes classe A da Internet já foram designadas e não temos como comprar mais redes desse tamanho. (E isso seria financeiramente impossível, porque custariam os “olhos da cara”.) Mas o porquê das 126 redes será explicado agora.

As várias classes de redes da Internet são identificadas pelo valor do seu primeiro octeto (para ser mais preciso, pelo valor dos quatro primeiros bits do primeiro octeto). Ou seja, só de olhar para um endereço IP válido na Internet, dá para saber a que classe de endereço ele pertence. Novamente, preste atenção, essa forma de endereçamento é apenas válida na Internet. Em uma rede local, as regras de endereçamento podem ser diferentes das praticadas na Internet, mesmo porque, em uma rede local privada, a IANA não tem competência nem autoridade para designar endereços ou faixas de endereços.

Vamos à forma de identificar as classes apenas analisando o primeiro octeto do endereço IP:

| <b>Classe</b> | <b>1<sup>o</sup><br/>octeto<br/>começa<br/>com<br/>(em<br/>binário)</b> | <b>1<sup>o</sup><br/>octeto<br/>pode ser<br/>(em<br/>decimal)</b> |
|---------------|-------------------------------------------------------------------------|-------------------------------------------------------------------|
| <b>A</b>      | <b>0</b>                                                                | <b>1      até<br/>126</b>                                         |
| <b>B</b>      | <b>10</b>                                                               | <b>128    até<br/>191</b>                                         |

|   |      |                |
|---|------|----------------|
| C | 110  | 192 até<br>223 |
| D | 1110 | 224 até<br>239 |
| E | 1111 | 240 até<br>254 |

Explicando melhor a relação entre os quatro primeiros bits do primeiro octeto e o valor do primeiro octeto.

Se o primeiro octeto (que é um número binário de 8 bits) começar com 0, é sinal de que ele pode ser 00000000 até 01111111 (ou seja, em decimal seria 0 até 127). Depois explico por que o primeiro octeto não pode ser 0 (zero) nem 127 na Internet. Sobra o quê? De 1 até 126.

Qualquer micro na Internet que tenha seu endereço IP com o primeiro octeto de valor contido dentro da faixa 1 até 126 pertence, automaticamente, a uma rede classe A.

Como exemplo, o micro 18.14.234.192 pertence a uma rede classe A.

Vamos continuar: se o primeiro octeto do endereço IP começar com 10 (em binário), as possibilidades seriam de 10000000 até 10111111 (em decimal seria de 128 até 191). Portanto, todo micro cujo endereço IP tenha como primeiro octeto algum valor contido entre 128 e 191 (inclusive eles) pertence automaticamente a uma rede classe B.

O endereço 173.14.254.9 pertence a uma rede classe B.

Continuando: qualquer endereço que começar com 110 (em binário) permite as seguintes variações: 11000000 (192) até 11011111 (223). Logo, qualquer endereço IP cujo primeiro octeto

esteja contido entre 192 e 223 faz parte da classe C.

Faz parte da classe C o micro 200.249.243.1.

A classe D é usada para endereços de multicast (não trataremos aqui) – são endereços usados para a criação de “grupos de computadores”. Ou seja, os endereços da classe D são usados não para identificar um micro em si (como os endereços das classes A, B e C), que costumamos chamar de endereços unicast (ou seja, para identificação única de um micro), mas os endereços de multicast servem para enviar uma mensagem a vários micros específicos em uma rede (um “grupo” específico de micros).

A classe E foi criada para que se tenha uma reserva de endereços para fins futuros. Talvez nunca venha a ser usada porque o IP, como conhecemos, está caindo em desuso. Dentro de alguns anos seremos obrigados, na Internet, a usar a *nova versão do IP (IPv6)* e a forma de endereçamento vai mudar radicalmente. (Veremos isso depois.)

“Ah João! Agora entendi por que só podem existir 126 redes classe A na Internet!”

Muito bem, leitor! É fácil entender isso.

Veja bem, se as redes classe A são, necessariamente, representadas apenas pelo primeiro octeto e este pode ser apenas de 1 até 126, existem apenas 126 redes classe A no mundo.

“E as restantes, João?”

As redes classe B foram designadas para serem usadas em ambientes onde o número de computadores é menor que 65.534 (lembre-se de que na classe B só podemos ter no máximo esse número de micros). São, ao todo, 16.384 redes classe B possíveis na Internet. Cada uma dessas redes pode ter, no máximo, 65.534 micros (hosts) – isso, claro, por causa da máscara das redes classe B (255.255.0.0).

São 16.384 redes possíveis (não sei se todas já foram vendidas pela IANA) porque o primeiro octeto pode variar de 128 até 191 (64 variações) e o segundo octeto (que também faz parte do identificador – ID – da rede) pode variar 256 vezes (de 0 a 255). Daí, 64 x 256 dá 16.384 possibilidades.

Isso significa que, se uma empresa tem o direito de comprar da IANA a rede 175.14.0.0 (lembre-se de que esse é o primeiro endereço daquela faixa classe B e é usado para identificar a rede), outra empresa teria o direito de comprar a rede 175.15.0.0. (O primeiro octeto é o mesmo, mas a rede é diferente, porque na máscara classe B, os dois primeiros octetos representam a rede.)

As redes classe C foram designadas para serem usadas em redes pequenas (com até 254 micros). São, ao todo, 2.097.152 redes classe C possíveis na Internet. Cada uma dessas redes pode ter, no máximo, 254 micros (hosts) – isso, claro, por causa da máscara das redes classe C (255.255.255.0).

O Brasil, basicamente, é composto de redes classe C ligadas à Internet (mesmo os nossos “provedores grandes”, como Terra e UOL, não compraram para classes A ou B por não poderem fazê-lo – escassez dessas redes para venda). Portanto, os provedores brasileiros compraram seus endereços diretamente da Embratel (que os comprou da IANA) – a maioria dos endereços IP do Brasil começa com o octeto 200 (a Embratel comprou inúmeras redes classe C começando com 200).

“Por que são mais de 2 milhões de possíveis redes classe C, João?”



Simples! O primeiro octeto das redes classe C pode ser qualquer valor entre 192 e 223 (incluindo ambos), o que dá 32 possibilidades. O segundo octeto e o terceiro octeto do endereço IP classe C também identificam a rede e, como variam de 0 a 255, permitem, cada um, 256 combinações possíveis.

Logo, temos  $32 \times 256 \times 256$  possibilidades de ID de rede classe C. Logo, são 2.097.152 redes classe C comercializáveis pela IANA na Internet. (Não, eu não sei se já foram todas vendidas, mas creio, sinceramente, que não!)

#### 8.11.5.1. Então, que tal um resumo rápido?

| Classe   | Primeiro Octeto                            | Nº de Redes |
|----------|--------------------------------------------|-------------|
| <b>A</b> | 00000001<br>(1) até<br>01111110<br>(126)   | 126         |
| <b>B</b> | 10000000<br>(128) até<br>10111111<br>(191) | 16.384      |
|          | 11000000<br>(192) até                      |             |

|          |                                            |        |
|----------|--------------------------------------------|--------|
| <b>C</b> | 11011111<br>(223)                          | 2.097. |
| <b>D</b> | 11100000<br>(224) até<br>11101111<br>(239) | *      |
| <b>E</b> | 11110000<br>(240) até<br>11111110<br>(254) | *      |

#### 8.11.5.2. Faixas de endereços reservados a redes privadas

A IANA é responsável pelos nomes e endereços usados para identificar os recursos na Internet, não nas redes privadas das empresas que se ligam à Internet (a menos que tais empresas queiram comprar endereços IP da IANA ou de algum de seus subsidiários).

Mas, pensando justamente em quem não quer comprar endereços junto a ela, a IANA reservou três faixas de endereços IP que podem ser usadas livremente em qualquer rede privada. (Rede cujos endereços IP não são válidos perante a Internet, ou seja, não foram comprados diretamente de órgãos como a IANA.)

Observe que essas três faixas denotam endereços IP que não podem ser usados livremente na Internet (ou seja, nenhum micro na Internet apresentará qualquer endereço IP dentro dessas faixas) porque eles simplesmente não são vendidos na Internet (nem pela IANA, nem pela ICANN) – esses endereços são usados para que uma empresa possa nomear seus micros em seu âmbito de rede privada apenas.

- **10.0.0.0 até 10.255.255.255:** é uma rede classe A que começa com o octeto 10 (esse octeto, claro, não poderá ser usado diretamente na Internet, mas apenas entre os micros de uma rede privativa). É usado quando a empresa precisa de um esquema de endereçamento

de rede para identificar mais de 65.534 computadores.

Os micros de uma rede com essa configuração poderiam ser: 10.23.245.35, 10.34.52.108, 10.2.2.2, 10.0.0.7 e assim por diante.

- **172.16.0.0 a 172.31.255.255:** é uma faixa de endereços originalmente localizada na classe B da IANA, mas que permite a existência de 16 redes classe B distintas (cada uma delas com capacidade para 65.534 computadores) – são as redes que vão dos octetos 172.16, 172.17, 172.18 e assim por diante, até 172.31.

- **192.168.0.0 até 192.168.255.255:** uma faixa de endereços reservados que permite a existência de 256 redes classe C distintas. Cada rede, claro, devido à máscara classe C, vai possibilitar 254 micros. São as redes 192.168.0, 192.168.1, 192.168.2 e assim sucessivamente até 192.168.255.

Por exemplo: 192.168.7.12 e 192.168.7.123 estão dentro da mesma rede (a rede 192.168.7.0). Já os micros 192.168.23.18 e 192.168.23.254 fazem parte de outra rede (a rede 192.168.23.0). Como esses endereços podem ser implantados naturalmente em uma rede privada (porque nunca serão usados diretamente na Internet), uma empresa pode ter várias sub-redes (divididas pelo terceiro octeto, neste caso).

Se uma empresa quiser que seus micros tenham endereços válidos diretamente na Internet (endereços globalmente únicos – ou seja, nenhum outro micro na Internet vai ter aqueles endereços), ela tem de comprá-los (registrá-los) em um órgão competente, como a IANA ou qualquer um de seus subsidiários. Os endereços das três faixas (blocos) mostradas neste tópico não são válidos perante a Internet, por isso, uma empresa poderá organizar seus micros internamente sem a necessidade de contatar nenhuma instituição (afinal, só para repetir, os endereços só serão válidos dentro daquela rede).

### 8.11.6. Endereços IP especiais

A IANA estipulou alguns endereços que não poderão ser usados em micros na Internet (nem em redes locais). Já vimos dois deles:

- **Endereço IP com ID do Host “tudo zero”:** também chamado de “*Host ID all zeroes*”, é o endereço da rede em questão. Lembre-se de que esse é o primeiro endereço IP de uma faixa analisada com uma determinada máscara e não pode ser usado para identificar micros porque identifica a rede em si (é “o nome da rede”).

Chamamos de ID do Host “tudo zero” porque simplesmente o primeiro endereço de uma faixa é o ID da rede (que não muda na rede inteira) seguido de octetos totalmente zero.

O micro 23.12.345.2, pertence a uma rede cuja máscara é 255.255.0.0. Qual é o endereço da rede? Simples: 23.12.0.0 (é o primeiro endereço da faixa que tem fixos os octetos 23.12). Ou se preferir explicar assim: colocou-se 0 em todo o ID do Host (parte do endereço IP que identifica uma máquina específica).

- **Endereço IP com ID do Host “tudo um”:** também conhecido como “*Host ID all ones*”.

Este é o endereço de broadcast de uma rede específica. É o último endereço daquela rede. Ele serve para enviarmos mensagens para todas as máquinas daquela rede específica.

Ou seja, se for necessário enviar uma mensagem a todos os micros da rede 23.12.0.0 (máscara 255.255.0.0), podemos utilizar o endereço 23.12.255.255.

“Certo, João. Então por que dizer ‘tudo um’ e não ‘tudo 255’?”

Porque o endereço IP é binário! E 23.12.255.255 é o mesmo que:

00010111.00001100.11111111.11111111

(Tudo “um” na parte do Host ID – os dois últimos octetos do endereço IP de máscara 255.255.0.0)

Os demais endereços IP reservados (proibidos para serem usados em micros) são:

- **0.0.0.0 (“all zeroes” ou “tudo zero”)**: é um endereço especial usado no momento em que uma estação faz o boot (inicialização) pela rede (solicitando o endereço IP que irá usar daquele momento em diante). Então, quando um micro não sabe qual é seu próprio endereço IP (no momento em que é ligado e vai pedir um para prosseguir), ele usa o endereço 0.0.0.0 como sua “assinatura” no primeiro pacote que ele envia para a rede.
- **1.1.1.1 (“all ones” ou “tudo um”)**: é um endereço que atua como broadcast local (qualquer pacote enviado para 1.1.1.1 será enviado a todos os micros da mesma rede em que o remetente estiver).

Observe que não é um broadcast para a rede que você quiser (feito no item “b” deste tópico), mas um broadcast para a mesma rede a que pertence o micro remetente. (Chamado também de broadcast limitado – porque não passa pelo roteador.) Neste caso os roteadores simplesmente não deixam pacotes endereçados a esse IP serem passados para outras redes.

- **127.x.x.x (Endereço de Loopback – autorretorno)**: qualquer endereço que tenha seu octeto inicial 127 (em binário 01111111) é considerado inválido para identificar micros porque esse endereço identifica a própria máquina em si. De uma forma bem simples, é uma maneira de um micro dizer “eu”. Ou seja, cada micro refere-se a si mesmo como sendo 127.x.x.x (normalmente, 127.0.0.1).

A explicação mais aprofundada (usando camadas) é que: quando o pacote é criado com o endereço de destino começando com 127 (pode ser 127.45.23.44, não importa), ao descer para a camada de enlace (camada 2), o quadro não é construído. O pacote simplesmente volta para a camada de inter-redes (camada 3), como se estivesse sendo recebido.

Ou seja, o protocolo IP envia o pacote para a camada 2, que, por sua vez, devolve-o ao protocolo IP, que “acredita” estar recebendo um pacote a ele direcionado.

Também é chamado de endereço do **Local Host**.

### 8.11.7. Analisando máscaras de sub-rede binárias

Vimos que as máscaras de sub-rede mais comuns são 255.0.0.0 (usada nas redes classe A), 255.255.0.0 (usada nas redes classe B) e 255.255.255.0 (classe C). Vimos também que quanto maior o ID da rede (octetos com 255), maior o número possível de redes e menor o número de micros por cada rede.

Claro que o inverso é fácil de deduzir: quanto menor o ID da rede, maior será o ID do host (a parte composta por octetos 0). Isso fará o número de redes possíveis diminuir, mas o número de hosts em cada rede aumentar.

Uma coisa que dá muita dor de cabeça para a maioria dos concurseiros (e para quem estuda redes por “hobby” ou na faculdade) é o cálculo com máscaras de sub-rede que usam números diferentes de 255 e 0 (como, por exemplo, a máscara 255.255.240.0).

“Isso cai em prova de concurso, João?”

Não é comum, especialmente se mencionarmos as provas da FCC e Cespe. As únicas bancas examinadoras que PODEM exigir isso (pelo histórico das provas recentes) são a ESAF e a FGV (especialmente esta última).

Vamos lá... Se você se deparar com uma máscara formada por octetos não completamente 1 (11111111 – ou 255) nem completamente 0 (00000000 – ou 0 simplesmente), a resolução terá de ser feita através da conversão do endereço IP para binário.

Atenção: se você não sabe converter um número decimal em binário, pare agora mesmo! Vá à parte de Aritmética Computacional (Parte 12 deste livro) e leia os tópicos que falam sobre isso (Conversão de Decimal para Qualquer Base e posteriormente Conversão de Qualquer Base para Decimal). Vamos lá. Eu espero. Você encontrará inclusive um tópico por lá que foi escrito especialmente sobre o endereço IP.

“Tudo bem, João. Podemos continuar. Já dá para fazer a conversão!”

OK, vamos começar com uma coisa simples (e que você já conhece). Temos um micro cujo endereço IP é 192.168.214.7 com a máscara de sub-rede 255.255.0.0 (Classe B). Qual parte é o ID da rede? Qual é o endereço da rede? Qual é o endereço de broadcast dessa rede? Quantos micros pode haver nessa rede?

|                             | <b>Octeto<br/>1</b> | <b>Octeto<br/>2</b> |
|-----------------------------|---------------------|---------------------|
| <b>End. IP<br/>do Micro</b> | 192                 | 168                 |
| <b>Máscara</b>              | 255                 | 255                 |
| <b>ID Rede</b>              | 192                 | 168                 |
| <b>End. IP<br/>da Rede</b>  | 192                 | 168                 |
| <b>End.</b>                 |                     |                     |

|                         |                      |     |
|-------------------------|----------------------|-----|
| Broadcast               | 192                  | 168 |
| Número de Hosts na Rede | $M = 256^K - 2 = 25$ |     |

(\*) M é o número de micros da rede / K é o número de octetos 0 da máscara.

Vamos analisar o mesmo exemplo em binário? Convertendo os octetos.

|                  | Octeto 1 | Octeto 2 |
|------------------|----------|----------|
| End. IP do Micro | 11000000 | 10101000 |
| Máscara          | 11111111 | 11111111 |

Podemos escrever simplesmente assim:

|              |                          |
|--------------|--------------------------|
| Endereço IP: | 110000001010100011010110 |
| Máscara:     | 111111111111111110000000 |

Agora vamos à regra de ouro: em um endereço IP visto de forma binária, é considerado ID da

rede aquela parte do endereço IP que coincide com os bits 1 da máscara. Logo, é considerado ID do host a parte do endereço IP que coincide com os bits 0 da máscara.

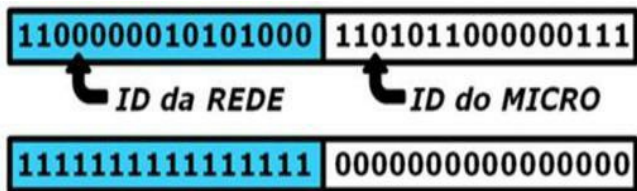


Figura 8.71 – Endereço IP e máscara binários.

Logo, dá para perceber que o ID da rede presente no endereço mostrado no exemplo é 1100000010101000. Com base na aquisição do ID da rede, dá para saber o endereço da rede e o endereço de broadcast daquela rede.

Para o endereço da rede, basta preencher o restante dos bits (ou seja, a parte que seria o ID do host) com 0 (zero) e, para o endereço de broadcast, basta completar o restante dos bits com 1. Simples assim. (Finalmente entendemos a história do “tudo um” e “tudo zero” no ID do Host.)

|                  | Octeto 1  | Octeto 2    |
|------------------|-----------|-------------|
| End. IP do Micro | 1 1000000 | 101 0110000 |
| Máscara          | 1 1111111 | 1 1111111   |
| ID da Rede       | 1 1000000 | 101 0110000 |
| End. IP          | 1 1000000 | 101 0110000 |

|                                 |                          |     |
|---------------------------------|--------------------------|-----|
| da Rede<br>End. de<br>Broadcast | 1 1000000                | 101 |
| Núm. de<br>Hosts da<br>Rede     | $M = 2^K - 2 = 2^{16} -$ |     |

(\*) A ideia é a mesma do cálculo de número de hosts na máscara decimal. Apenas com uma diferença.

Cálculo do número de hosts (M) em uma máscara binária:  $M = 2^K - 2$

Sabemos que  $M$  é o número que queremos descobrir (número máximo de micros que aquela rede suporta que sejam ligados). Mas na máscara decimal, K representa o número de octetos com valor 0 (haverá, no máximo, 3 deles). Na máscara binária, o ***K representa o número de bits 0*** (algarismos binários). Observe que a base que será elevada a K mudou (na máscara decimal, a base é 256, aqui na máscara binária, a base é 2).

Bom, para máscaras classe A, B, e C o raciocínio é o mesmo de seus cálculos em decimal, mas para máscaras que vão além do 255 e 0, o “buraco é mais embaixo”... Nesses casos, só dá para fazer por meio dos cálculos binários. Vamos a um exemplo.

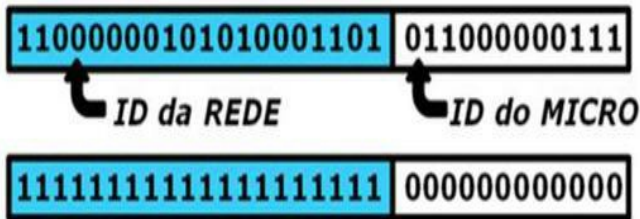
Temos um micro cujo endereço IP é 192.168.214.7 com a máscara de sub-rede 255.255.240.0. Qual parte é o ID da rede? Qual é o endereço da rede? Qual é o endereço de broadcast dessa rede? Quantos micros pode haver nela?

“João, essa máscara é classe A, B ou C?”

Nenhuma das três. É uma máscara de sub-rede no sentido literal da expressão. Aqui já não classificamos como uma rede classe A, classe B ou classe C. Temos uma definição de uma rede classe B dividida em sub-redes. (É o que acontece quando temos máscaras que usam mais que simplesmente 255 e 0.)

Aqui já entramos no que conhecemos como CIDR (Classless Inter-Domain Routing – Roteamento entre Domínios sem Classe) – que é justamente a técnica de endereçamento que usa máscaras diferentes das classes A, B e C. Veja uma representação gráfica do micro e da máscara descritos neste exemplo que estamos analisando.





*Figura 8.72 – Os 20 primeiros bits representam o ID da rede.*

Para responder às perguntas, temos de, impreterivelmente, fazer a análise com base nos endereços binários em questão. Então vamos lá! Note que o terceiro octeto não é 00000000 (que em decimal seria 0), mas 11110000 (que vale 240). Note ainda que é no octeto 3 que a separação entre ID da rede e ID do host acontece (daí eu ter “separado” com um espaço o que é 1 do que é 0 na máscara).

|                                        | Octeto 1 |  |
|----------------------------------------|----------|--|
| End. IP do<br>Micro<br>(192.168.214.7) | 11000000 |  |
| Máscara<br>(255.255.240.0)             | 11111111 |  |
| ID da Rede                             | 11000000 |  |

|                            |                 |
|----------------------------|-----------------|
| End. IP da Rede<br>(*)     | 11000000        |
| End. de<br>Broadcast (*)   | 11000000        |
| Número de<br>Hosts da Rede | $M = 2^K - 2 =$ |

(\*) Convertendo os endereços IP para a notação decimal separada por pontos, o endereço IP da rede seria 192.168.208.0 e o endereço de broadcast daquela rede seria 192.168.223.255. (Para não fugir à regra, esses são, respectivamente, o primeiro e o último endereços de uma rede com a máscara 255.255.240.0 e da qual faz parte o micro 192.168.214.7.)

Quanto ao número de hosts (micros) possíveis naquela rede, é fácil. Já havíamos visto isso. Na máscara 255.255.240.0 (em binário 11111111111111111111000000000000), há 20 bits 1 (representando o prefixo da rede) e há 12 bits 0 (representando o espaço do endereço que irá variar de micro para micro naquela rede). É só fazer o cálculo considerando os 12 bits 0.

|           |          |       |       |
|-----------|----------|-------|-------|
| $M = 2^K$ | $2^{12}$ | 4.096 | 4.094 |
| - 2;      | - 2;     | - 2;  | hosts |

Só para mostrar uma coisa muito interessante: em vez de se referir a essa rede como a rede “192.168.208.0 com máscara 255.255.240.0”, costuma-se usar outra denominação ainda mais fácil de lembrar: 192.168.208.0/20. Essa notação com / (barra) é a mais usada em endereços CIDR.

“Fantástico, João! É mesmo! É mais fácil de escrever. Mas, o que significa o 20?”

Simples, essa notação significa que a rede se chama 192.168.208.0 (esse é seu primeiro endereço IP) e que existem 20 bits 1 (bits que representam o ID da rede) na máscara de sub-rede (ou seja, ela é 11111111111111111111000000000000 ou 255.255.240.0).

Vamos a mais uma... Tomando-se a rede 159.48.0.0/13, quais micros a seguir estariam inseridos nesta rede? E quantos micros podem existir nesta rede?

Micro A – 159.50.12.78

Micro B – 159.120.204.1

Micro C – 159.55.36.139

Em primeiro lugar, vamos entender que, para saber se um micro faz parte de uma rede, seu endereço IP tem de ter o mesmo ID de rede que o endereço IP da rede. Ou seja, o endereço IP do micro A (159.50.12.78), por exemplo, tem de ter o mesmo ID da rede que o endereço 159.48.0.0/13 (que é o endereço da rede).

“Como fazer isso, João? Como descobrir se eles têm o mesmo ID da rede?”

Simples. Vamos começar com o endereço da rede em si (temos de descobrir qual é a máscara e qual é o endereço da rede em binário). Só assim descobriremos o ID da rede (a parte do endereço IP que será idêntica para todos os micros daquela rede).

IP da rede: 159.48.0.0/13: são 13 bits 1 no início da máscara! Portanto, já se tem a máscara em binário (11111111.11111000.00000000.00000000)

|                          |                                                   |
|--------------------------|---------------------------------------------------|
| IP da rede (em binário): | <b><i>10011111.00110000.00000000.00000000</i></b> |
| Máscara (em binário):    | <b><i>11111111.11111000.00000000.00000000</i></b> |

Observe que a parte que está em destaque no IP da rede é exatamente a parte que coincide com os bits 1 da máscara, logo, é ela o ID da rede que estamos procurando. (Lembre-se de que o ID da rede é a parte do endereço IP que nunca mudará para todos os micros daquela mesma rede!)

Portanto, o ID da rede em questão é 10011111.00110 (pega todo o primeiro octeto e os cinco primeiros bits do segundo octeto). Agora é só analisar quais micros (dentre os três apresentados) possuem o mesmo ID de rede em seus endereços IP.

|                         |  |
|-------------------------|--|
| <b><i>Micro A –</i></b> |  |
|-------------------------|--|

|                                  |                               |
|----------------------------------|-------------------------------|
| <b>159.50.12.78</b>              |                               |
| <b>IP do Micro (em binário):</b> | <b>10011111.00110010.0000</b> |
| <b>Máscara (em binário):</b>     | <b>11111111.11111000.0000</b> |

Levando em consideração apenas os 13 primeiros bits do endereço IP (não precisaríamos nem mesmo desenhar a máscara, não é?), o ID da rede do micro A é idêntico ao ID da rede do endereço dado para a nossa rede de exemplo. Logo, o micro A faz parte da rede 159.48.0.0/13.

**Micro B – 159.120.204.1**

**IP do Micro (em binário): 10011111.01111000.11001100.00000001**

Não precisamos nem usar a máscara, não é mesmo? (Porque a função da máscara é apenas identificar que parte é ID da rede – onde houver 1.) Como sabemos que os 13 primeiros bits são 1, basta considerar os 13 primeiros bits do endereço IP para achar o ID da rede.

No caso do micro B, o ID da rede dele é diferente do ID da rede do nosso exemplo! Logo, o **micro B não faz parte da rede 159.48.0.0/13.**

**Micro C – 159.55.36.139**

**IP do micro (em binário): 10011111.00110101.00100100.10001011**

Analisando, novamente, apenas os 13 primeiros bits do endereço IP do micro C convertido para binário, temos que o ID da rede dele é idêntico (10011111.00110) ao ID da rede que estamos analisando. Logo, ele também pertence à rede do nosso exemplo.

Já sabemos que os micros A e C fazem parte da rede. Parte da nossa questão já foi resolvida. Agora resta saber quantos micros podem ser instalados naquela rede. Isso, meus amigos, é muito fácil! Não é necessário nem saber qual é o IP da rede. Basta que saibamos o que vem depois da “p”.

Portanto, para descobrir a quantidade de micros na rede 159.48.0.0/13, só precisamos do 13! Vamos à nossa linha de raciocínio.

O número 13 indica que há 13 bits 1 no início da máscara de sub-rede. Esses 13 bits 1 indicam

que parte do endereço IP identifica o ID da rede (os 13 primeiros bits, lógico). Como o endereço IP é formado por 32 bits, restam 19 para o campo do ID do host. Portanto, a máscara é composta de 13 bits 1 seguidos de 19 bits 0.

Exatamente assim: 11111111111100000000000000000000

Se já temos a quantidade de bits 0 (19 no exemplo), já podemos aplicar a fórmula!

$$M = 2K - 2 = 219 - 2 = 524.288 - 2 = 524.286 \text{ hosts}$$

“Ah, João, entendi! Para saber o número de hosts de uma rede, é só ter em mãos o número de bits 0 da máscara! Então, a parte 159.48.0.0 nem era necessária. Como você bem disse, só o número que vem depois da barra (/), no caso o 13 é que é necessário para conseguirmos essa informação!”

Precisamente, caro leitor. Exatamente isso!

### 8.11.8. Como o meu micro recebe os parâmetros IP?

Um computador pode receber as informações necessárias para conexão na rede IP (os chamados parâmetros IP – ou seja, seu endereço próprio, o endereço do gateway padrão e a máscara de sub-rede) basicamente de duas maneiras.

Normalmente citamos como sendo o fornecimento do “endereço IP” apenas, mas o raciocínio serve para os demais parâmetros (máscara e endereço do gateway):

**Endereço IP Fixo:** é fornecido ao computador pelo administrador da rede (ou pelo administrador do próprio computador, se este tiver autonomia para tanto). Esse endereço é configurado diretamente dentro das propriedades do computador (no sistema operacional), e este computador sempre vai apresentar esse endereço (mesmo que seja desligado, o computador vai “acordar” depois ainda possuindo esse número).

O endereço fixo é usado quando se conhece o endereço de todos os outros computadores, ou seja, quando se tem o controle da rede inteira (daí o fato de ser o administrador a atribuir os endereços IP fixos). O endereço IP fixo também é usado em servidores (computadores que detêm informações na Internet) para que estes sempre estejam no mesmo endereço, prontos para responder às requisições dos demais computadores.

**Endereço IP Dinâmico:** é usado em todas as conexões domésticas à Internet. Nesse caso, o endereço IP (e os demais parâmetros) é fornecido ao computador no momento em que este se conecta à rede, e “esquecido” quando o computador é desligado da rede. Cada vez que um computador se conecta à rede, ele recebe um endereço IP, que normalmente é diferente dos anteriores (dentro da faixa definida de possibilidades daquela rede – ou sub-rede).

Um computador recebe o endereço IP dinâmico de um servidor que usa um protocolo chamado DHCP (Dynamic Host Configuration Protocol – Protocolo de Configuração Dinâmica de Host). O protocolo DHCP é, em suma, o responsável pela atribuição automática de endereços IP aos computadores na rede (vamos falar um pouco mais sobre ele adiante).

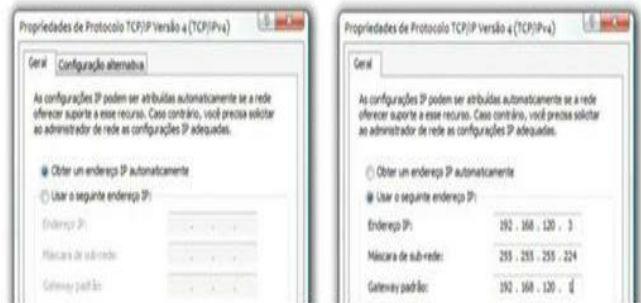


Figura 8.73 – Endereço IP dinâmico (à esquerda) e endereço IP fixo.

Com isso, terminamos a análise do IP atual. Sim, IP atual. Vamos mergulhar agora no que vem por aí!

#### 8.11.9. IPv6 – nova forma de endereçamento na Internet

Tudo o que estudamos agora (endereço IP, máscara de sub-rede, classes) está para se tornar inútil (na verdade, já deveria ter se tornado, segundo os agendamentos, neste último ano de 2012).

Toda a Internet deve mudar seu estilo de endereçamento para uma versão mais nova do protocolo IP, a chamada versão 6 (IPv6). Cumpre salientar, caro leitor, que tudo que se viu até agora neste livro trata da versão 4 (IPv4).

Vamos analisar um pouco, mesmo que superficialmente, o protocolo IP na versão 6.

Em primeiro lugar, o IPv6 promove uma mudança radical (demais) no endereço IP. Ou seja, os computadores passarão a ter endereços diferentes daqueles com que estamos acostumados.

“Mas, João, justo agora que entendi tudo sobre máscara e endereço IP e tudo mais. Já tava gostando da coisa.”

É, caro leitor... Só que o endereço IP como conhecemos existe desde a década de 1970. Já estava na hora de mudar. Imagine, por exemplo, que o endereço IP atual tem 32 bits e isso permite  $2^{32}$  combinações diferentes de endereços. São, ao todo, mais de 4 bilhões de endereços.

Mesmo parecendo ser muita coisa, estes 4 bilhões já são considerados insuficientes, seja porque muitos desses endereços estão inutilizados (redes classe A e B subutilizadas desperdiçando endereços), seja porque há muito mais equipamentos (além de computadores) ligados à Internet hoje em dia (celulares, TVs, videogames, geladeiras, todos eles precisam de endereços IP para se conectar).

Então, já havia uma preocupação latente em quando os endereços IP da Internet passariam a ser considerados escassos. Eles já são considerados escassos há muito tempo, e a mudança começou há pelo menos uma década.

A criação do IPv6 não é de agora, mas a obrigatoriedade da migração é coisa recente. Logo, logo, teremos todos de “migrar” (na verdade, não somos nós, são os nossos provedores de acesso) para usar IPv6.

“Certo, João! Mas o que há de tão diferente ou especial nesse IPv6?”

Para começo de conversa, o endereço IP na versão 6 não tem 32 bits. Tem “miseráveis” 128 bits! Ele ficou simplesmente quatro vezes maior!

011110000010110111010010000101111010000001011000111101110110000111011000101100

Dá para se lembrar de algo assim? Claro que não? Se fôssemos representar o endereço IPv6 em notação decimal com pontos, seria algo assim:

120.45.210.23.208.44.123.176.236.89.10.45.23.69.45.127

Mas claro que não o representamos dessa forma. Também não dá para se lembrar desse endereço, não é mesmo?! Os endereços IPv6 são escritos como oito grupos de dígitos hexadecimais separados por dois-pontos.

782D:D217:D02C:7BB0:EC59:0A2D:1745:2D7F

“João, pelo amor de Deus! De onde surgiu isso aí?”

**Conversão dos números para a base hexadecimal**, caro leitor! Vá ao capítulo final deste livro para ficar mais familiarizado com a Conversão de Binário para Hexadecimal e a Conversão de Hexadecimal para Binário.

Mas é claro que as provas poderão apresentar endereços IPv6 com alguma(s) de suas “formas” otimizadas. Vamos às regras que permitem escrever os endereços IPv6 mais facilmente.

a) Os zeros à esquerda de cada grupo podem ser omitidos (como em qualquer número, na verdade). Desta forma, o endereço

5002:0023:0AB4:120A:B0BA:B234:0001:0453

poderá ser escrito assim:

5002:23:AB4:120A:B0BA:B234:1:453

b) Grupos compostos por apenas 0 (zero) podem ser omitidos, sendo substituídos por uma dupla de dois-pontos. Mesmo que haja vários grupos seguidos compostos apenas de zeros. Ou seja, o endereço:

2301:0000:0000:0000:0012:023A:00AA:04BC

poderá ser escrito assim:

2301::12:23A:AA:4BC

Claro que já considere também omitir os 0 à esquerda (regra “a”)

Outra informação importante: não é permitido fazer uso de “::” mais de uma vez em um endereço, sob pena de ambiguidade na interpretação, segundo se verifica a seguir:

2A2B::13A::20C3

poderia significar 2A2B:0000:0000:013A:0000:0000:0000:20C3

mas também poderia significar 2A2B:0000:013A:0000:0000:0000:0000:20C3

Se você analisar bem, caro leitor, com o IPv6 temos agora endereços suficientes para 2<sup>128</sup>

estações diferentes. É mais ou menos equivalente a  $3,4 \times 10^{38}$  micros. Esse número é maior que o Número de Avogadro (e eu nem mais sei o que é o Número de Avogadro).

### 8.11.9.1. Endereços IPv6 especiais

Alguns endereços IPv6 são especiais (assim como existem em IPv4). O primeiro a ser visto é o endereço de loopback (autorretorno), que representa a própria máquina que está enviando o pacote.

Na estrutura do IPv4, este endereço é qualquer um que comece com 127 (primeiro octeto), como 127.0.0.1 (mais usado). No IPv6 é diferente (claro!): o endereço de loopback é apenas 1 (sim... 127 números 0 seguidos de um único 1). Seria assim (em hexadecimal):

**0000:0000:0000:0000:0000:0000:0000:0001**

Ou, se preferir (e sei que vai), assim:

**::1**

Há um endereço IPv6 que não será atribuído a nenhum micro: o endereço “**tudo zero**” (todos os 128 bits 0). Esse endereço é usado para indicar, em qualquer troca de pacotes, a ausência de endereço válido (endereço nulo, não funcional). Ele é representado assim:

**::**

Quando qualquer tráfego (pacotes) de uma rede IPv6 precisar passar por alguma rede IPv4, ele precisa “se adaptar” ao jeito antigo dessa rede. Essa adaptação consiste, entre outras coisas, em apresentar, no pacote, um endereço IPv4 compatível com aquela rede anciã por onde o pacote está passando.

Para encapsular um pacote IPv6 em uma rede IPv4, usa-se um endereço IPv4 escrito da mesma maneira como aprendemos (notação decimal separada por pontos) precedido de :: (dois sinais de dois-pontos).

**::192.168.23.45**

Quando um computador que só usa endereços IPv4 é colocado para trabalhar em uma rede que usa endereços IPv6, a rede faz uma “gambiarra” para permitir que aquele micro ancião possa fazer parte daquela rede novinha em folha. Essa “gambiarra” é feita adicionando-se 16 bits 1 antes do endereço IPv4 daquela estação. Representamos isso em notação hexadecimal (separada por dois-pontos) e decimal (separada por pontos). É, as duas. Veja só:

**::FFFF:192.168.23.45**

É interessante notar que o número 192.168.23.45 pode ser representado em hexadecimal também (através de dois grupos de quatro dígitos), portanto o endereço anterior poderia ser escrito também como:

**::FFFF:C0A8:172D**

### 8.11.9.2. Tipos de endereços IPv6

Há basicamente três tipos de endereços nas redes que se baseiam em IPv6:

- **Unicast:** endereço que identifica uma única máquina (é um endereço IPv6 que está associado a um único computador);
- **Anycast:** endereço IPv6 que identifica um grupo de computadores, mas a mensagem



enviada a um endereço anycast será entregue a qualquer um dos micros daquele grupo (não a todos). Normalmente escolhe-se o “mais próximo” computador como destinatário. O termo “any” representa bem a ideia de “qualquer um”.

- **Multicast:** endereços IPv6 que representam várias máquinas (um grupo) e fazem com que a mensagem seja entregue a todas elas. Quando um pacote é enviado a um endereço de multicast, ele será entregue a todas as máquinas daquele grupo!

É fácil reconhecer endereços de **multicast** porque seus oito primeiros bits são 1, fazendo sua representação normal, em hexadecimal, iniciar com os dígitos FF:

**FF45::1349:D3A4:90A2**

Informação importante: não há endereço de broadcast nas redes IPv6! Ou seja, não existe nenhum endereço que atue como “enviar para todos os micros da rede”.

### 8.11.9.3. ID de rede e ID de host no IPv6?

Essa preocupação também foi palco de discussão entre os que desenvolveram o IPv6. O endereço IPv6 não tem um “ID de rede” (não com esse nome). O endereço IPv6 possui um prefixo da sub-rede.

O prefixo da sub-rede é o conjunto dos primeiros n bits do endereço (esse “n” é variável). Para saber quantos bits do endereço IPv6 são reservados para identificar a rede da qual o micro faz parte, usa-se a notação da barra (/), como no IPv4.

Então, a notação

**F4D3:23AB:A120:4AEF::2AA9/60**

indica que os primeiros 60 (sim, 60!) bits daquele endereço representam o prefixo da sub-rede na qual o micro em questão está ligado. Os demais 68 bits do endereço são a parte do endereço que identifica o computador em si.

“João, não tem mais coisa sobre IP não, tem? Pelo amor de Deus!”

Ter tem, caro leitor... Mas acho que para o nosso foco de estudo já foi alcançado. Na verdade, acho que vimos até coisa demais, mas que vai propiciar a você capacidade de responder a qualquer questão da Esaf (ou FGV, ou Cespe, ou qualquer uma!)

Vamos estudar agora os demais protocolos da camada 3 (inter-redes).

### 8.11.10. Protocolo ICMP

O protocolo ICMP (Internet Control Messaging Protocol – Protocolo de Mensagens de Controle de Inter-Redes) trabalha justamente com algo com que o IP não trabalha: a detecção de erros nos pacotes que trafegam pela Internet.

Funciona mais ou menos assim: quando um roteador recebe um pacote IP contendo dados, ele analisa aquele pacote, a fim de descobrir se há algum problema. Se não houver, ótimo, o pacote é encaminhado à próxima rede; mas se houver algum problema naquele pacote, o roteador em questão constrói, por meio do protocolo ICMP, uma mensagem de erro (ou mensagem de controle) e a envia em um pacote IP ao emissor daquele pacote defeituoso, pedindo que se tomem as providências necessárias (como o reenvio).

Note que a mensagem ICMP é encapsulada em (colocada dentro de) um pacote IP normal. A única diferença é que os roteadores vão saber que esse pacote IP não contém dados, mas sim

uma mensagem de controle (mensagem ICMP).

As mensagens ICMP que podem ser trocadas entre os vários dispositivos de rede (roteadores e/ou estações) são várias. Cada uma delas tem um nome e um número (chamado tipo). Não acho necessário decorá-las, mas aqui vão algumas delas:

- **Echo Request (8) e Echo Reply (0):** essas duas mensagens são trocadas quando um emissor deseja saber se um receptor está “vivo” e ativo (respondendo). O micro A envia um Echo Request para o micro B, que deverá, se estiver funcionando corretamente, enviar de volta para A uma mensagem chamada Echo Reply.

O comando Ping (um comando do Linux e do Windows que serve como “testador” de comunicação) se baseia no uso de mensagens Echo.

“João, e esses (8) e (0) são os tipos respectivos das mensagens?”

Sim, caro leitor. Esses são os números que identificam os tipos das mensagens na documentação oficial que rege a Internet.

- **Timestamp Request (13) e Timestamp Reply (14):** são semelhantes às mensagens Echo, só que com uma diferença: na Timestamp é registrada a hora exata em que aconteceu o Request e o Reply.

- **Destination Unreachable (3):** é a mensagem enviada por um roteador (ou host) ao emissor de um pacote IP original com o intuito de informar que aquele pacote não poderá ser entregue ao destino. O motivo da não entrega é especificado dentro do pacote que contém a mensagem.

- **Time Exceeded (11):** é a mensagem ICMP enviada por um roteador ao emissor de um pacote IP quando o campo TTL (Time-To-Live) desse pacote atinge 0. Ou seja, a mensagem Time Exceeded é o “telegrama” que informa ao emissor que “aquele pacote morreu”!

- **Source Quench (4):** é a sinalização que um roteador envia ao emissor do pacote para informar que aquele pacote foi descartado por causa de congestionamento da Internet. Mais precisamente, a questão é com aquele roteador, que deve estar com a memória totalmente esgotada para processar qualquer pacote extra. Então o roteador simplesmente “diz” aos demais: “Já chega! Não aguento mais!”.

- **Redirect (5):** essa mensagem demonstra muita “humildade” por parte do roteador. Quando um roteador “percebe” que não é o ideal para que um pacote IP de dados atinja seu destino, ele simplesmente envia uma mensagem Redirect para o emissor do pacote de dados original.

Nessa mensagem, o roteador assume sua “incompetência” para encaminhar pacotes para aquele endereço, “pede desculpas” e informa o endereço IP do roteador que ele “crê” ser o mais indicado para processar os próximos pacotes de dados para aquele mesmo destino. Que retidão e que senso de autocritica, hein?

- **Parameter Problem (12):** essa mensagem é enviada ao emissor de um pacote quando se encontra algum problema na própria estrutura do pacote (algum campo com defeito ou mal transmitido).

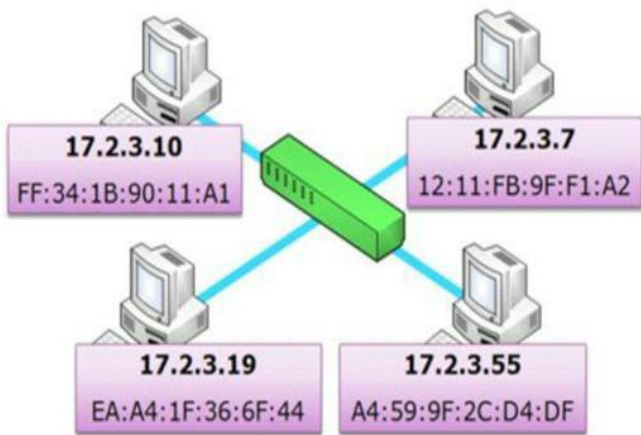
Essa análise é feita com base no checksum (soma de verificação) do cabeçalho do pacote IP original. Se um roteador fizer o cálculo do checksum e notar que ele “não bate” com o valor

contido no cabeçalho daquele pacote, ele deduz que houve problemas com o próprio conteúdo do cabeçalho do pacote e imediatamente dispara uma mensagem tipo 12 para o emissor!

#### 8.11.11. Protocolo ARP

É usado para permitir a associação de endereços IP (endereços lógicos) com endereços da camada de enlace. (Endereços MAC, também conhecidos como endereços de hardware, que são os endereços das placas de rede, como vimos anteriormente.)

O ARP (Address Resolution Protocol – Protocolo de Resolução de Endereços) é usado para associar um endereço IP a um endereço de hardware. Isso quer dizer: quando um endereço IP é fornecido para a entrega de um determinado datagrama, o computador que detém aquele endereço é localizado por meio do ARP, que lê o endereço IP e aponta qual o endereço MAC do computador que o possui. Vamos analisar isso melhor na figura a seguir.



*Figura 8.74 – Quatro estações e seus endereços IP e MAC.*

Imagine, caro leitor, que o micro 17.2.3.10 queira enviar um pacote a um certo micro 17.2.3.7. Fácil imaginar isso, não? Mas tem um problema: qualquer pacote a ser enviado tem de se transformar em um quadro (sim, lembra disso?) antes de ser enviado pela estrutura física da rede.

E para ser transformado em um quadro, deve-se saber o endereço MAC (endereço físico) desse destinatário. Ai é que está! O micro remetente sabe apenas o IP do destinatário. (Isso é suficiente para construir um pacote, mas não para colocar esse pacote num quadro e enviá-lo pela rede.)

É aí que entra o ARP. Para descobrir o endereço MAC do destino, o micro remetente consulta sua tabela ARP interna (chamada Cache ARP ou Tabela ARP) para ver se existe alguma associação que permita descobrir qual é o MAC do micro cujo IP é o 17.2.3.7.

Caso não encontre, será enviada uma mensagem em broadcast naquela rede (ao endereço MAC de broadcast, que é FF:FF:FF:FF:FF:FF ou “tudo um”) perguntando “Qual de vocês é o dono do IP 17.2.3.7?”. Essa mensagem é chamada **ARP Request**.

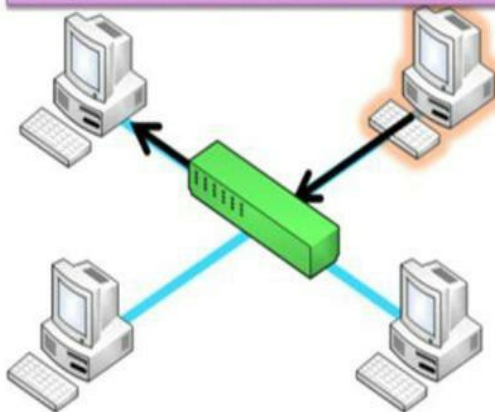


Figura 8.75 – Uma mensagem ARP Request (Requisição ARP).

A mensagem chegará a todos os micros daquela rede. Os micros analisarão a pergunta e consultarão os endereços IP de suas próprias placas de rede. Caso alguém ache o endereço IP em questão, mandará como resposta o endereço MAC da placa de rede que está associada ao IP originalmente consultado. Isso é o **ARP Reply** (resposta ARP).

**De: 12:11:FB:9F:F1:A2    Para: FF:34:1B:90:11:A1**

**Eu!!! Eu tenho 17.2.3.7!**



*Figura 8.76 – Uma mensagem ARP Reply (Resposta ARP).*

#### **8.11.12. Protocolo RARP**

O “irmão inútil” do ARP é o protocolo RARP (Reverse ARP – ARP Reverso).

“Ô, João, maldade... Por que ‘inútil’?”

Porque esse protocolo não é mais usado, caro leitor. Ele servia para o exato inverso do ARP, ou seja, com base na informação de um endereço MAC inicial, o RARP conseguia descobrir o endereço IP da máquina.

“E quando isso era usado?”

Notei o seu espanto. Percebeu que é uma coisa completamente sem sentido? Mas era usado para quando um computador com placa de rede e sem disco rígido (estação diskless) queria buscar o sistema operacional em um servidor e precisava de um IP. Como não há disco rígido naquela estação, ela não poderia guardar o IP dentro de si.

Como a estação tinha placa de rede, ela tinha endereço MAC. Restava apenas descobrir o IP para associar àquele MAC. Portanto aquela estação sem disco enviava um quadro RARP com

seu endereço MAC pedindo “Pelo amor de Deus, alguém me dê um endereço IP!!!”, então um bondoso e paciente servidor RARP dava um endereço IP àquela estação que tanto o queria.

O protocolo RARP foi substituído por dois sucessores: o BOOTP e o DHCP. (Este último é usado na Internet atualmente, pois se tornou padrão desde 1993).

Vamos subir uma camada para conhecer os protocolos contidos na Camada de Transporte.

## 8.12. Protocolos de transporte

A camada de transporte do modelo TCP/IP é composta, originalmente, por apenas dois protocolos, cuja responsabilidade, como citado anteriormente, é estabelecer uma conexão fim a fim entre os dois hosts (computadores) envolvidos na comunicação.

“João, dá para lembrar o porquê de fim a fim?”

Claro! Os protocolos da camada de transporte não se preocupam como a mensagem vai trafegar pela Internet (o IP se preocupa com isso) nem tampouco com a transmissão da mensagem dentro de uma mesma rede (o protocolo da camada de interface de rede faz isso). Em vez desses dois motivos de preocupação, os protocolos de transporte simplesmente se preocupam com a “quebra” da mensagem em vários segmentos (na origem) e a reunificação de tais segmentos no destino.

É responsabilidade dos protocolos da camada de transporte criar mecanismos (incluir informações no cabeçalho dos segmentos) que permitam que a reunificação aconteça de forma perfeita e, com isso, que a mensagem chegue ao seu destino inteira (ou quase).

Os protocolos que formam essa camada são:

- **TCP**
- **UDP**

### 8.12.1. Protocolo TCP

O protocolo TCP (Transmission Control Protocol – Protocolo de Controle de Transmissão) é um protocolo de transporte orientado a conexão. Seu funcionamento é bem simples e ao mesmo tempo bem estruturado para garantir a transmissão dos pacotes entre os computadores envolvidos na comunicação.

“João, o que é ‘ser orientado a conexão’?”

Em poucas palavras, quer dizer que o protocolo TCP faz com que o emissor só comece a transmitir seus dados se tiver certeza de que o receptor está pronto para ouvi-los. Ou seja, toda a transmissão se orienta pelo estabelecimento de uma conexão prévia entre os dois envolvidos. Não há transmissão sem que haja uma conexão estabelecida entre eles.

Por ser orientado a conexão, o TCP traz uma série de características que são consequência disso:

- **É confiável:** garante a entrega de todos os dados no destino sem defeito ou perda.
- **Garante a sequência dos segmentos:** os segmentos que saem do emissor são numerados e reunidos na mesma ordem no micro de destino.
- **Reconhecimento:** o receptor envia um segmento de confirmação (reconhecimento) para cada segmento de dados que receber, informando ao emissor que ele já poderá transmitir o

próximo segmento da sequência.

- **Retransmissão:** se um segmento se perder (por causa de problemas de transmissão nas demais camadas), o TCP do receptor solicitará ao TCP do emissor o reenvio do segmento faltoso.

- **Deteção de duplicidade:** o TCP reconhece se um segmento chegou em duplicidade no receptor e automaticamente descarta o segmento duplicado.

- **Controle de fluxo:** o emissor não vai enviar mais segmentos do que a quantidade que o receptor for capaz de processar (mesmo porque o emissor só transmitirá quando o receptor informar que ele pode fazê-lo).

- **Controle de congestionamento:** o TCP ajusta-se automaticamente às quedas de desempenho da rede provocadas por congestionamento (nos roteadores e servidores, por exemplo).

- **Estabelece sessões:** o TCP trabalha por meio do estabelecimento de sessões de comunicação, em que várias transmissões são feitas em bloco e consideradas parte de uma sessão só.

- **Troca informações de estado (status):** os dois hosts ligados em TCP trocam entre si constantemente informações de apresentam o status da conexão entre eles.

- **Baixa velocidade:** devido à grande quantidade de informações, recursos e itens que garantem a integridade das transmissões via TCP, é fácil deduzir que o protocolo TCP não é tão rápido quanto seu “irmão inconsequente”.

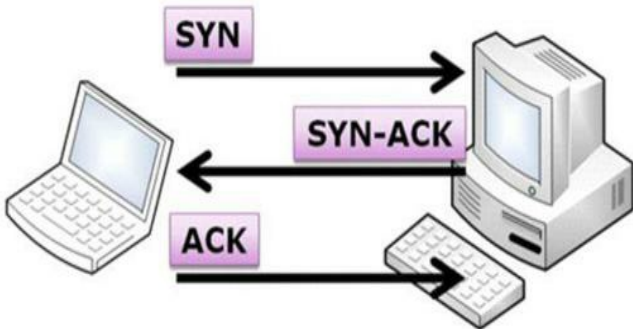
#### 8.12.1.1. Estabelecimento de conexão TCP (aperto de mãos em três vias)

Antes de transmitir qualquer pacote de dados, os dois computadores que pretendem usar o protocolo TCP como protocolo de transporte na comunicação entre eles devem estabelecer uma conexão.

Fica mais fácil de entender se usássemos uma comparação com a linha telefônica: “Antes de duas pessoas que moram distante uma da outra transmitirem qualquer fofoca entre si, elas têm de estabelecer uma conexão telefônica entre elas.”

Ou seja, antes de começar com “Menina... Você soube da mais nova?”, um dos envolvidos tem de discar o número do outro e o outro tem de atender.

No caso de dois computadores com TCP, eles trocam entre si três segmentos simples antes de poderem transmitir qualquer dado. Esses três segmentos não contêm dados em si, mas as informações necessárias para o início daquela conexão.



*Figura 8.77 – Three-way Handshake (Aperto de mãos em três vias).*

É simples:

1. O micro que inicia o processo (o cliente) envia um segmento vazio, contendo apenas uma informação no cabeçalho do TCP (o bit SYN com valor 1) – esse segmento é chamado segmento SYN.
2. O micro com o qual o primeiro quer se conectar (neste caso é o servidor) recebe o segmento SYN e constrói e envia de volta ao cliente o pacote com os bits SYN e ACK definidos com valor 1. Este é chamado segmento SYN-ACK.
3. Para finalizar o estabelecimento da conexão, o micro cliente recebe o SYN-ACK e devolve um segmento com o bit SYN em 0 e o bit ACK em 1. Esse segmento, como já deve ter dado para deduzir, é o segmento ACK.

“João, o que são os ‘bits SYN e ACK’?”

São partes do cabeçalho do segmento TCP (campos, como aqueles que vimos anteriormente quando estudamos o formato do pacote IP). Não creio ser necessário saber o formato do segmento TCP, não. Depois da troca bem-sucedida dos três segmentos de estabelecimento de conexão, os dois micros já têm condições de trocar entre si segmentos TCP contendo dados reais (dados dos aplicativos envolvidos).

### 8.12.2. Protocolo UDP

O protocolo UDP (User Datagram Protocol – Protocolo de Datagrama de Usuário) é um protocolo de transporte sem conexão que fornece uma entrega rápida, mas não confiável, dos pacotes. Esse protocolo é uma opção em relação ao TCP e usado em menos casos.

Por ser um protocolo não confiável, ele não fornece o controle de fluxo necessário, nem



tampouco exige uma confirmação do receptor, o que pode fazer com que a perda de um pacote aconteça sem a devida correção. Por isso ele é usado em aplicações nas quais a velocidade é mais importante que a integridade dos dados (como vídeos e música pela Internet).

Pelo fato de não exigir confirmação do receptor quanto à chegada dos pacotes, o protocolo UDP não sobrecarrega a rede tanto quanto o TCP (afinal, cada confirmação de recebimento é um pacote sendo transmitido, não é?), mas também por causa disso, não é confiável.

“João, e quem é que usa essa ‘porcaria’, se ela não é confiável?”

O serviço de DNS, por exemplo, que veremos depois, usa UDP como protocolo de transporte, porque deseja velocidade. O protocolo TFTP (FTP Trivial) também usa UDP. Serviços que permitem ouvir músicas e assistir a vídeos diretamente pela Internet também foram desenvolvidos para usar o UDP em vez do TCP.

### 8.12.3. Resumo TCP versus UDP

Segue um pequeno resumo que poderá ajudar quando esses conceitos forem exigidos em uma prova qualquer:

| UDP                                                                | TCP                                                                      |
|--------------------------------------------------------------------|--------------------------------------------------------------------------|
| Serviço sem conexão; nenhuma sessão é estabelecida entre os hosts. | Serviço orientado por conexão; uma sessão é estabelecida entre os hosts. |
| UDP não garante nem                                                | TCP garante a entrega                                                    |

confirma a  
entrega dos  
dados, nem  
organiza em  
sequência os  
mesmos.

através do  
uso de  
confirmações  
e entrega  
sequenciada  
dos dados.

Os programas  
que usam UDP  
são  
responsáveis  
por oferecer a  
confiabilidade  
necessária ao  
transporte de  
dados.

Os  
programas  
que usam  
TCP têm  
garantia de  
transporte  
confiável de  
dados dada  
pelo próprio  
protocolo.

UDP é rápido,  
necessita de

TCP é mais  
lento,

baixa  
sobrecarga.

necessita de  
maior  
sobrecarga.

“Ei, João... E nós podemos escolher se usamos o TCP ou o UDP?”

Não, caro leitor, a menos que sejamos os desenvolvedores (programadores) de algum aplicativo. Pois, como usuários, apenas utilizamos programas que já estão prontos, inclusive com suas definições precisas de qual será o protocolo de transporte que irão utilizar.

#### 8.12.4. Portas

Quando estamos usando a Internet, seja navegando em uma página, mandando um e-mail ou batendo um papo (ou, quem sabe, tudo junto), o protocolo de transporte usado é, na grande maioria dos casos, o TCP. Mesmo que vários processos/programas usem o mesmo protocolo de transporte, como saber que um determinado pacote que chegou é destinado àquele ou a esse programa? Fácil: Através de um número de identificação chamado *porta* (*port*).

Cada serviço que usamos (e-mail, Web, bate-papo, FTP etc.) usa, necessariamente, um protocolo de aplicação diferente (SMTP, HTTP, IRC, FTP etc.), que veremos a seguir. Cada um desses protocolos usa o protocolo TCP (ou UDP) para garantir o transporte das mensagens que enviam para outros computadores e, quando o fazem, informam ao TCP (ou ao UDP) o número da porta que vão usar para realizar a transferência.

Cada micro abre uma porta específica para uma comunicação específica. Aquela porta aberta servirá para que aquele computador identifique aquela comunicação enquanto ela estiver acontecendo.

“João, explica aí essa estória de ‘abrir uma porta’. Como é isso?”

Caro leitor, é a coisa mais simples do mundo. Uma porta é um número (um campo) no cabeçalho do segmento do protocolo de transporte. Uma porta é apenas um identificador. Quando um programa (aplicativo cliente, por exemplo) pede para estabelecer uma comunicação com um aplicativo servidor, ele o faz na forma de pacotes (como já vimos antes) e estes vão conter, além dos dados já vistos, como endereço IP de origem e destino, a porta através da qual se deseja fazer aquela transferência.

Um computador servidor pode abrir para si uma porta X (digamos 80), a fim de estabelecer um processo de comunicação com um computador cliente que abriu a porta Y (digamos 12.456) para a mesma comunicação (aliás, isso acontece o tempo todo).

A porta é como o nome da pessoa a quem uma carta se destina. Exemplo: ao escrever para Jorge, da família Santos, residente na Avenida dos Pinhais, 450, você está definindo não somente o endereço do destinatário (esse endereço é o mesmo para qualquer pessoa que more com Jorge), mas sim especificando para quem lá dentro a carta vai. O carteiro só tem responsabilidade de entregar a carta na casa, depois disso, a família separa as cartas destinadas a

cada um dos seus integrantes.

Se você entendeu, é assim: Av. dos Pinhais, 450 é o endereço do computador (como, por exemplo, 200.245.150.13), e Jorge Santos é o endereço da porta para quem o pacote (carta) deve ser entregue. Essa porta está, necessariamente, associada a um determinado serviço usado na Internet.

Existem, ao todo, 65.536 portas disponíveis, da porta 1 à porta 65.536. Mas as mais usadas vão somente da 1 até a 1.024 (algumas raras estão acima disso).

“Meu Deus, João! Tenho de conhecer todas as 1.024 comuns?”

Claro que não, caro leitor, basta umas 500 (brincadeira!). Vamos apresentar os números das principais portas à medida que apresentarmos os protocolos de aplicação.

#### **8.12.5. Socket**

Dá-se o nome de socket (soquete) ao conjunto de informações formado por endereço IP do destinatário, protocolo e porta de comunicação. Ou seja, quando um computador estabelece uma conexão com o endereço 200.234.15.129 pela porta 80 do protocolo TCP, simplesmente estabeleceu um “socket”. Toda comunicação na pilha TCP/IP é realizada por meio de sockets.

Ou seja, quando você envia um e-mail, acessa uma página, baixa um arquivo ou realiza qualquer outra operação na Internet, deve estabelecer um socket com o seu interlocutor (o computador do outro lado).

Um socket é, em poucas palavras, uma maneira de um computador identificar, de forma única, sem ambiguidades ou dúvidas, um determinado processo de comunicação que esteja estabelecido.

Ou seja, um socket é um “registro” de uma comunicação, uma “assinatura” que permite ao nosso computador, mesmo se comunicando com vários micros ao mesmo tempo, saber a quem pertence aquele pacote específico que chegou a ele (porque chegou pelo socket específico).

### **8.13. Protocolos de aplicação**

São os protocolos descritos da última camada do modelo, que entram em contato com o usuário, permitindo que este possa se comunicar com os demais componentes do seu computador e enviar suas mensagens pela rede até outros computadores. Os protocolos dessa camada estão associados diretamente aos principais serviços usados pelo usuário na rede: e-mail, Web, bate-papo etc. Os principais protocolos de aplicação são:

#### **8.13.1. SMTP**

SMTP (Simple Mail Transfer Protocol – Protocolo de Transferência Simples de Correio) é o protocolo usado para o envio de mensagens de correio eletrônico (e-mail). Esse protocolo usa a porta 25 do protocolo TCP.

Atenção: esse protocolo é usado no ato do envio do correio eletrônico. Não só no envio que acontece entre usuário remetente e servidor de correio, mas também entre servidor de envio e servidor de recebimento.

#### **8.13.2. POP**

POP (Post Office Protocol – Protocolo de Agência de Correio) é usado para realizar o recebimento das mensagens de correio eletrônico. Com esse protocolo, as mensagens armazenadas na caixa postal do usuário são trazidas para o computador do usuário e retiradas do servidor (a rigor, visto que se pode selecionar que as mensagens fiquem em cópia no servidor de e-mails). Esse protocolo usa a porta 110 do protocolo TCP. Atualmente encontra-se em sua terceira versão, daí o nome **POP3**.

### 8.13.3. IMAP

IMAP (Internet Message Access Protocol – Protocolo de Acesso a Mensagens na Internet) é usado em opção ao POP porque facilita o acesso aos dados nas caixas postais sem a necessidade de “baixá-los” para o computador cliente. Através do IMAP, é possível realizar um acesso on-line aos dados na caixa postal localizada no servidor sem que isso signifique trazer as mensagens ao micro do usuário.

É uma opção interessante para aqueles que pegam suas mensagens de e-mail de vários computadores diferentes. Todo acesso é feito através de aplicações que acessam a caixa postal, leem seu conteúdo e o mostram ao usuário. As caixas postais dos “webmails” (Gmail, Yahoo, Hotmail entre outros) usam o IMAP, pois os usuários têm acesso através de uma página Web, que mostra as mensagens e dá direitos de lê-las, apagá-las, respondê-las e tudo mais. O protocolo IMAP usa a porta 143.

### 8.13.4. HTTP

HTTP (Hyper Text Transfer Protocol – Protocolo de Transferência de Hiper Texto) é o protocolo usado para realizar a transferência das páginas Web para nossos computadores. O HTTP é usado para trazer o conteúdo das páginas (documentos feitos com a linguagem HTML) para nossos programas navegadores (Browsers). O protocolo HTTP utiliza a porta 80 do protocolo de transporte TCP.

Sim, é bom que se saiba que o HTTP não é seguro, portanto...

Há uma variação do HTTP, que se chama **HTTPS (HTTP Seguro)**, e é usado para realizar o acesso a páginas com transferência criptografada de dados (através de um algoritmo de criptografia chamado SSL). Esse protocolo é comumente usado nos acessos aos sites de bancos e lojas virtuais onde se informam números de cartão de crédito, por exemplo.

O HTTPS é, na verdade, a junção do HTTP, usado para transferir páginas, com o SSL (Secure Socket Layer), um protocolo de segurança, criado para fornecer criptografia aos protocolos que naturalmente não fazem uso dela (falaremos sobre ele mais adiante).

O protocolo HTTPS não é 100% seguro, ou seja, ele não evita completamente a ameaça de interceptação das mensagens entre usuário e site, mas oferece um nível de segurança que minimiza bastante esse risco. O protocolo HTTPS é usado sobre a porta 443.

### 8.13.5. FTP

FTP (File Transfer Protocol – Protocolo de Transferência de Arquivos) é usado para realizar a transferência de arquivos entre dois computadores através da Internet. O protocolo FTP exige o

estabelecimento de uma sessão (com o uso de login e senha).

O protocolo FTP utiliza duas portas no protocolo TCP: a porta 21 (da qual muitos se lembram) é usada para os comandos da conexão, como os que solicitam a listagem de diretórios, a cópia de arquivos e o apagamento deles etc., porém, a transferência dos dados propriamente ditos acontece pela porta TCP 20. Portanto, para a conclusão da transferência de um arquivo pelo FTP, são usadas duas conexões (sockets) diferentes.

Um parente próximo do protocolo FTP é o TFTP (FTP Trivial), que realiza a transferência de arquivos através do protocolo UDP e não do TCP, como seu irmão mais conhecido, o que permite uma transferência de arquivos com mais velocidade e sem uma série de recursos que o FTP oferece. O TFTP usa a porta 69.

Além de transferir arquivos, o protocolo FTP permite que o usuário realize uma gama enorme de operações com o micro a que se conectou. O FTP permite que pastas e arquivos sejam criados, excluídos, renomeados, movidos e copiados no servidor. Ou seja, basicamente tudo aquilo que se pode fazer no seu micro por meio do Windows Explorer é possível fazer em um servidor remoto por meio de FTP.

Claro que vale lembrar que o micro a ser controlado deve ter um programa aplicativo servidor de FTP atuando e que o login e a senha do usuário deem a ele o direito de fazer tais operações.

### **8.13.6. Telnet**

TELNET (Terminal Emulator – Emulador de Terminal) é um protocolo que realiza a conexão entre dois computadores para que um deles “finja” ser terminal do outro. Isso significa que qualquer comando executado no computador “terminal” será realizado, na verdade, no computador-alvo: o servidor.

Esse sistema era muito utilizado nos primórdios das redes de computadores, quando não se tinha dinheiro para fazer redes com computadores individuais interligados. A estrutura de “rede” normalmente consistia em um único computador central (o “console” ou “mainframe”), e os demais “computadores” eram apenas teclados e monitores ligados a esses (chamados terminais ou “terminais burros”). Todos os comandos executados nos terminais são realizados na CPU e na RAM do console.

Ou seja, um terminal não é um micro. Um terminal é apenas um “braço” de um computador. Não tem RAM, CPU, HD etc. Um terminal é apenas um teclado e um monitor.

Na verdade, os dois computadores envolvidos pela conexão do Telnet são microcomputadores, como os nossos; apenas um deles “finge” ser um terminal (o cliente), enquanto o outro “finge” ser um console central (o servidor). Todos os comandos digitados no teclado do “terminal” são realizados, na verdade, pela CPU e pela memória do computador central. O Telnet utiliza a porta 23 do protocolo TCP.

### **8.13.7. NNTP**

NNTP (Network News Transfer Protocol – Protocolo de Transferência de Notícias em Rede) é usado no serviço conhecido como News (Notícias), que reúne vários usuários em torno de newsgroups (grupos de notícias). Esse serviço é bastante semelhante a um serviço conhecido como Fórum (como o do site [www.forumconcurseiros.com](http://www.forumconcurseiros.com), que todos vocês, concurseiros,

conhecem). O protocolo NNTP utiliza a porta 119 do protocolo TCP.

#### 8.13.8. DNS

DNS (Domain Name Service – Serviço de Nome de Domínio) é um serviço usado para realizar a tradução dos nomes de domínios (URLs) em endereços IP. Ou seja, quando digitamos, em nosso navegador, “[www.euvoupassar.com.br](http://www.euvoupassar.com.br)”, esse endereço é enviado para um servidor que trabalha com o protocolo DNS, e que, por sua vez, devolve ao computador que requisitou o endereço IP associado ao domínio desejado. O serviço de DNS utiliza a porta 53 no protocolo UDP!

É o DNS que estabelece a estrutura hierárquica e organizada dos domínios como conhecemos atualmente na Internet (veremos mais adiante, no capítulo de Internet).

#### 8.13.9. DHCP

DHCP (Dynamic Host Configuration Protocol – Protocolo de Configuração Dinâmica de Estação) é um protocolo que fornece as informações IP necessárias para as estações poderem se ligar na rede.

Funciona de forma semelhante ao RARP: uma estação, ao se conectar à rede, envia uma solicitação a todos os micros da rede (essa mensagem é chamada de **DHCP Discover** – ou Descobrimiento DHCP). Na verdade, sem muito romantismo, é um pacote simplesmente enviado ao endereço de broadcast da rede.

A mensagem poderá chegar a vários servidores DHCP (computadores com capacidade de fornecer as informações IP às demais estações), visto que nessa rede pode haver vários servidores. Os servidores DHCP então enviam um pacote chamado **DHCP Offer** (ou Oferecimento DHCP), que contém um endereço IP disponível para aquele micro.

Sim, aquele micro que gritou pedindo um endereço IP poderá receber vários como resposta. É aí que ele faz a seleção! Esse micro escolhe um dos IP oferecidos e responde ao servidor que ofereceu endereço IP escolhido com uma mensagem chamada **DHCP Request** (Solicitação DHCP) que visa requisitar a confirmação da configuração que aquele servidor havia oferecido.

Por fim, o servidor responde ao micro requisitante com uma mensagem **DHCP Ack** (Confirmação Positiva DHCP), e o vínculo está estabelecido. (Ou seja, aquele micro, daquele momento em diante, passa a ser conhecido pelo endereço IP que o servidor lhe forneceu.)

#### 8.13.10. SNMP

SNMP (Simple Network Management Protocol – Protocolo de Gerenciamento Simples de Rede) é um protocolo que permite o gerenciamento da situação dos nós da rede. O SNMP não está preso ao conjunto TCP/IP, e pode ser usado para controlar qualquer tipo de equipamento de rede como roteadores, servidores, estações, pontos de acesso etc. desde que estes possuam suporte a esse protocolo.

Através do SNMP, podemos enviar comandos a vários tipos de equipamentos de redes para que eles se desliguem, ou reiniciem, ou realizem essa ou aquela tarefa. É um protocolo que permite o “controle remoto” de vários dispositivos da rede.

### 8.13.11. RTP e RTCP

O RTP (Real Time Protocol – Protocolo de Tempo Real) e o RTCP (Real-Time Control Protocol – Protocolo de Controle em Tempo Real) são usados para serviços que transferem grandes fluxos de dados em tempo real (ou seja, enquanto remetente e destinatário estão realmente se comunicando).

Alguns dos serviços que fazem uso desses dois protocolos são a transferência de música e vídeo pela Internet e o VoIP (Voz sobre IP) – que é a “telefonia” pela Internet.

Os protocolos da pilha TCP/IP são os mais usados da atualidade porque, é óbvio, são os protocolos usados na Internet (a maior conexão entre redes do mundo). Esse padrão foi estabelecido como sendo o padrão de protocolos usados nesse ambiente ainda quando a Internet era apenas uma pequena conexão entre universidades americanas.

Mas outros protocolos existem e são citados em concursos públicos. Esses protocolos serão mostrados agora.

### 8.14. Outros protocolos conhecidos

Dentre os protocolos não pertencentes ao conjunto TCP/IP, podemos citar alguns poucos que já interessaram aos órgãos “fazedores” de provas:

- **Netbeui:** Protocolo criado pela IBM para redes locais de computadores. Esse protocolo admite até 255 computadores em uma rede. Mas sua característica mais forte é que ele não é roteável.

Ser roteável significa que um protocolo pode ser lido por roteadores, e, portanto, pode ser usado em estruturas inter-redes (ou seja, em ligações entre redes). Já que essa não é uma das características do Netbeui, podemos concluir que ele não pode ser usado em Inter-redes (consequentemente, na própria Internet).

Onde usamos o Netbeui? Nas “redes Windows”. Ou seja, nas redes locais em que só se utiliza o sistema operacional Windows. O sistema Windows tem como principal protocolo de redes locais o Netbeui. Mas uma rede de computadores locais com Windows pode utilizar o Netbeui concomitantemente ao TCP/IP, o que permite que a referida LAN possa se conectar com a Internet (por causa do TCP/IP, não do Netbeui).

- **IPX/SPX:** É um conjunto de protocolos (assim como o TCP/IP) usado em redes de computadores Netware, da empresa Novell. As redes Netware são, na verdade, redes de computadores cujo servidor utiliza um sistema operacional chamado Netware, desenvolvido pela empresa Novell.

As redes Novell eram muito comuns, mas com o advento do Windows NT e seus sucessores, bem como do Linux como sistema operacional de servidores, o sistema Netware e a própria Novell vêm, gradativamente, perdendo espaço.

O IPX é um protocolo roteável localizado na camada de rede e é equivalente ao IP na pilha TCP/IP. O SPX é um protocolo da camada de transporte, equivalente ao TCP na pilha TCP/IP.

### 8.15. Considerações finais

Bem, chegamos ao fim de mais um assunto em nosso livro. Espero que tenham gostado da



parte de redes de computadores (embora eu saiba que é, em muitos casos, um assunto chato). O assunto que vocês acabaram de enfrentar é, na maior parte das vezes, exigido pela Esaf e pela FGV em nível bastante técnico. As demais elaboradoras de provas não “gostam” tanto de redes de computadores.

## 8.16. Questões de Redes de Computadores

1. O administrador da rede informou que o servidor SMTP que atende à empresa não está funcionando. Para os usuários que utilizam esse servidor, isso significa que, enquanto o problema persistir:
  - a) o único serviço prejudicado será o de recebimento de e-mail;
  - b) o envio e o recebimento de e-mail devem ser feitos pelo servidor POP3 e, consequentemente, esses dois procedimentos ficarão um pouco lentos;
  - c) os serviços de recebimento e envio de e-mail foram paralisados;
  - d) o único serviço prejudicado será o de envio de e-mail;
  - e) não será possível navegar na Internet.
2. TCP/IP é o nome que se dá ao conjunto de protocolos utilizados pela Internet. Esse conjunto de protocolos foi desenvolvido para permitir aos computadores compartilharem recursos em uma rede. Todo o conjunto de protocolos inclui padrões que especificam os detalhes de como conectar computadores, assim como também convenções para interconectar redes e rotear o tráfego. Oficialmente, esse conjunto de protocolos é chamado protocolo Internet TCP/IP, geralmente referenciado só como TCP/IP, devido a seus dois protocolos mais importantes. Considerando o modelo de referência OSI para o conjunto de protocolos TCP/IP, encontram-se dois protocolos: um deles oferecendo serviços sem conexão e o outro oferecendo serviços orientados a conexão. Esses dois protocolos localizados na camada de transporte são:
  - a) SMTP e POP3;
  - b) FTP e UDP;
  - c) TCP e http;
  - d) FTP e Telnet;
  - e) UDP e TCP.
3. Analise as seguintes afirmações, relacionadas aos conceitos básicos de redes de computadores, seus componentes, protocolos, topologias e servidores.
  - I. No modelo OSI, a camada de aplicação é responsável pelo endereçamento dos pacotes, convertendo endereços lógicos em endereços físicos, de forma que os pacotes consigam chegar corretamente ao destino. Essa camada permite que duas aplicações em computadores diferentes estabeleçam uma sessão de comunicação. Nessa sessão, essas aplicações definem como será feita a transmissão de dados e coloca marcações nos dados que estão sendo transmitidos.
  - II. O SMTP permite que um usuário, utilizando uma máquina A, estabeleça uma sessão interativa com uma máquina B na rede. A partir dessa sessão, todas as teclas pressionadas na máquina A são repassadas para a máquina B, como se o usuário tivesse um terminal ligado diretamente a ela.
  - III. O DNS é particularmente importante para o sistema de correio eletrônico. Nele são definidos registros que identificam a máquina que manipula as correspondências relativas a um determinado domínio.

IV. O FTP permite que um usuário em um computador transfira, renomeie ou remova arquivos remotos.

Indique a opção que contenha todas as afirmações verdadeiras.

- a) I e II.
- b) II e III.
- c) III e IV.
- d) I e III.
- e) II e IV.

4. Analise as seguintes afirmações, relativas a redes de computadores.

I. O endereço físico gravado em uma memória ROM dentro de uma placa de rede é denominado endereço MAC. Ele deve ser alterado de acordo com a máscara de sub-rede da LAN onde a placa será utilizada.

II. O protocolo NetBEUI é roteável e deve ser utilizado em LANs que não têm acesso direto à Internet.

III. O IPX/SPX é o protocolo proprietário criado pela Novell. O IPX opera na camada de rede e é o equivalente do IP no protocolo TCP/IP.

IV. O protocolo UDP não verifica se o pacote de dados chegou ao seu destino.

Indique a opção que contenha todas as afirmações verdadeiras.

- a) I e II.
- b) II e III.
- c) III e IV.
- d) I e III.
- e) II e IV.

5. Analise as seguintes afirmações relacionadas a conceitos básicos de redes de computadores.

I. No roteamento dinâmico utilizado pelos hubs e switches, as tabelas de roteamento refletem dinamicamente as modificações na topologia da rede. As tabelas são atualizadas a partir de informações trocadas entre estes dispositivos.

II. O endereço usado para identificar uma sub-rede, denominado máscara de sub-rede, deve ser composto por bytes completos. Dessa forma, em uma LAN, as três máscaras de sub-rede possíveis são: 255.255.255.0, 255.255.0.0 e 255.0.0.0.

III. Alguns endereços IP são reservados, não podendo ser utilizados para identificar as placas de interface de rede em um computador. Um desses endereços, o 127.0.0.0, identifica a própria máquina.

IV. Um ARP traduz um endereço IP para o endereço MAC correspondente. Quando o endereço MAC associado ao um endereço IP não é conhecido, o ARP envia uma mensagem de consulta para o endereço de broadcast. Cada máquina na rede recebe a mensagem e verifica se o endereço IP consultado pertence a uma de suas placas e, em caso afirmativo, responde informando o endereço MAC equivalente.

Indique a opção que contenha todas as afirmações verdadeiras.

- a) I e II;
- b) II e III;

- c) III e IV;
- d) I e III;
- e) II e IV.

6. Os switches são dispositivos:

- a) capazes de estabelecer a comunicação de computadores distantes entre si e até mesmo com protocolos de comunicação diferentes;
- b) utilizados por uma tecnologia de rede desenvolvida pela IBM chamada Token Ring, cujo princípio de operação é a comunicação em forma de circuito fechado;
- c) que têm a função de transferir os pacotes de um segmento para todos os demais, não fazendo qualquer tipo de seleção ou endereçamento;
- d) semelhantes a hubs, mas não repetem o mesmo pacote para todas as portas. Cada pacote é dirigido para o dispositivo de destino, evitando colisões e excesso de tráfego;
- e) da estrutura de nível mais alto em uma rede composta por várias sub-redes. O switch é composto por linhas de conexão de alta velocidade, que se conectam às linhas de menor velocidade.

### 9.1. O que é a Internet?

Essa eu sei, João! A Internet é a maior rede de computadores do mundo!”

Errado, caro leitor! Errado mesmo!

A Internet é a maior ligação entre redes de computadores do mundo. A Internet não é uma rede só: são várias. As redes que formam a Internet são interligadas por meio de roteadores, os equipamentos que vimos anteriormente para essa finalidade (interligar redes).

Se você não acredita, pense bem: sendo usuário do provedor XYZ, você pode mandar um e-mail para um funcionário da empresa WKG? *Claro que sim!* E, isso significa que as duas empresas mencionadas têm uma rede só de computadores? Uma rede única? *Claro que não!*

Logo, só existe uma explicação: a rede da empresa XYZ está interligada à rede da empresa WKG, assim como às demais empresas que, juntas, formam a Internet. Lembre-se de que Internet significa Inter-net (ou Inter-redes) – ou seja, ligação “entre redes”.

### 9.2. Como a Internet nasceu?

Esse é o tipo de informação de que você não vai precisar para fazer a prova, mas preferi mostrar para que você pudesse entender melhor a Internet.

Bem, na década de 1960 começaram os primeiros projetos no intuito de obter uma ligação entre computadores em longas distâncias. Essa ligação começou no MIT (Instituto de Tecnologia de Massachussets) e foi encomendada pela DARPA (Agência de Pesquisas e Projetos Avançados do Departamento de Defesa dos Estados Unidos).

Foram desenvolvidas, então, as principais características da *ARPAnet*, a rede que interligaria por alguns anos os principais computadores acadêmicos e militares dos Estados Unidos.

Durante a década de 1970, muitos padrões usados hoje foram desenvolvidos, como os protocolos do conjunto TCP/IP, o serviço de e-mail e o FTP, entre outros. Mas foi no início da década de 1980 que a rede ARPAnet se dividiu em algumas “sub-redes” e recebeu o nome “oficial” de Internet. Essa descentralização foi aumentada no decorrer dos anos seguintes, visto que a Internet se tornou o que é hoje: uma conexão mundial que conta com cerca de 1 bilhão de computadores conectados.

### 9.3. E a Internet aqui?

No Brasil, a Internet chegou com atraso: no final da década de 1980. O maior contribuinte para esse advento foi a *RNP* (Rede Nacional de Ensino e Pesquisa), criada pelo Ministério da Ciência e Tecnologia, que visava à interligação do ambiente acadêmico brasileiro em uma estrutura inter-redes. (Ou seja, essa rede serviria para ligar as universidades brasileiras tanto entre si, quanto com as inter-redes internacionais.)

Em 1995, a Internet se “abriu” comercialmente no Brasil, fazendo a RNP reestruturar seus modelos de serviço. A iniciativa privada, claro, deu sua contribuição para a solidificação das

conexões da Internet no Brasil. Várias empresas montaram seus **backbones** no Brasil, que, hoje, é um dos países mais “conectados” do planeta.

Em tempo: backbone significa “espinha dorsal”. É um termo usado para definir a estrutura principal de uma ligação inter-redes. Um backbone é constituído de todas as ligações (cabos, satélites) e equipamentos (roteadores, servidores, modems) que são encontrados no “meio” das redes. Se um backbone é a “espinha dorsal”, nós, os usuários, somos os “terminais”, os “dedos” da Internet. Veja, na figura a seguir, como está a RNP hoje (a figura mostra o backbone da RNP).

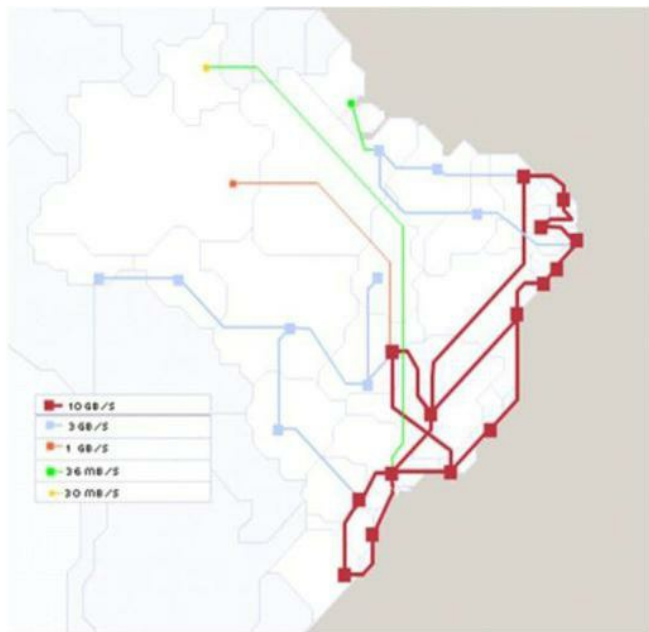


Figura 9.1 – Backbone da RNP (do site [www.rnp.br](http://www.rnp.br)).

### 9.3.1. Internet 2

A Internet 2 é um consórcio de várias universidades americanas (e parceiros), voltado para pesquisas avançadas em novos sistemas que transmitam dados em altíssima velocidade (planejam 100 Gbps).

Atualmente, diversas universidades, empresas e centros de pesquisa desenvolvem materiais e tecnologias para transmissão de dados com alto desempenho. A Internet 2 é a versão “não pública” da Internet. Não temos acesso a seus dados, apenas às explicações sobre sua forma e seu objetivo.

Muito provavelmente, as tecnologias que hoje são projetadas e testadas na Internet 2 virão a se tornar populares na “nossa” Internet.

### 9.4. Conectando-se à Internet

Costumamos nos conectar à grande rede através de um intermediário (uma empresa que está no “meio do caminho”), conhecido como **ISP** (Provedor de Serviços de Internet) ou **Provedor de Acesso** ou simplesmente **Provedor**.

Um provedor é uma empresa que se mantém conectada à estrutura da Internet constantemente e “repassa” esse acesso aos usuários, quase como um “cambista” que compra vários ingressos e os revende. Nós, usuários domésticos de Internet, e as empresas de pequeno e médio porte nos conectamos à Internet por meio de um provedor.



*Figura 9.2 – Conexão com a Internet.*

As grandes empresas (especialmente as de telecomunicações) eventualmente podem se conectar, também, por meio de provedores de acesso, mas normalmente se conectam diretamente ao backbone da rede, sendo, assim, suas próprias provedoras (as empresas de telecomunicações são parte do backbone).

Existem várias formas através das quais os usuários podem se conectar à Internet, que diferem entre si em velocidades, equipamentos e sistemas de comunicação utilizados. Aqui vão os métodos de conexão mais comuns com a Internet. (Na verdade, são sistemas de comunicação usados para diversos fins, mas hoje em dia, a Internet é esse fim.)

#### **9.4.1. Linha telefônica (Dial-Up)**

A conexão dial-up (termo que pode ser entendido como “por meio de discagem telefônica”) se dá por intermédio de uma linha telefônica convencional com o uso de um equipamento conhecido como **Modem** (modem telefônico, para ser mais exato).

Na conexão dial-up, o computador do usuário se conecta a um modem telefônico que, por sua vez, se liga à linha telefônica convencional. A ligação é realizada estabelecendo uma ligação telefônica mesmo (ou seja, consumindo pulsos telefônicos). O computador do provedor atenderá



a ligação e permitirá o acesso depois de comprovada a autenticidade do usuário.

Uma das características mais inconvenientes dessa forma de conexão é a velocidade: a taxa máxima de transferência desse sistema é de **56 Kbps** (56 Kilobits por segundo), que é o limite do modem. Outro ponto importante: enquanto se está conectado, os pulsos telefônicos são contabilizados, pois se trata de uma ligação telefônica local convencional. O equipamento usado nesse sistema de conexão é o modem telefônico (hoje fabricado sob as normas V.90 e V.92).

Nas ligações dial-up, é comum o uso de dois protocolos SLIP e PPP. O PPP (Protocolo Ponto a Ponto) é mais recente e usado pela maioria dos provedores de acesso atualmente.

“Ei, se é conexão com a Internet, não deveria usar o TCP/IP, João?”

E usa, caro leitor! O PPP e o antigo SLIP são protocolos da camada de enlace (camada 2), ou seja, estão relacionados intimamente com a forma física da conexão; portanto, trabalham abaixo do IP e do TCP, que podem ser usados sobre aqueles.

Lembre-se: assim como existe o CSMA/CD na camada de enlace das redes Ethernet, existem os protocolos de rede telefônica (PPP e SLIP), que funcionam para estabelecer a ligação entre os dois pontos da conexão telefônica (modem do micro do usuário e modem do provedor).

Detalhe importante: esta forma de conexão já está caindo em desuso, sendo utilizada, apenas, quando não há outra opção naquela localidade! Já é possível realizar conexão de Internet em casa por meio de diversos outros métodos mais rápidos!

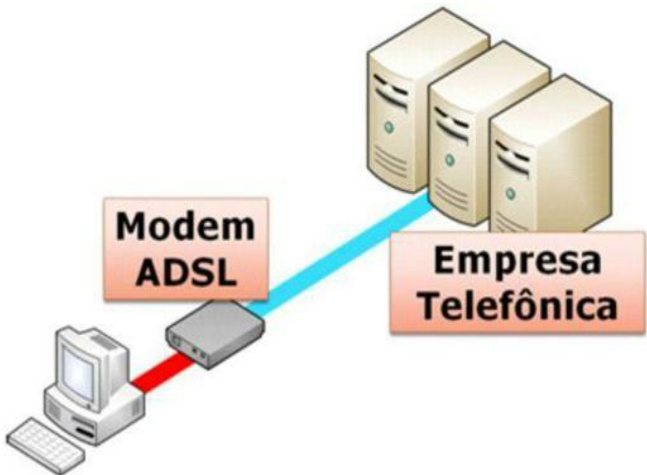
#### 9.4.2. ADSL

O serviço de **ADSL** (*Asymmetric Subscriber Digital Line – Linha de Assinante Assimétrica Digital*) é o nome técnico dado aos sistemas de acesso em banda larga (alta velocidade) oferecidos pelas empresas de telefonia fixa. O produto Velox, da Oi, o Speedy, da Telefônica, e a Internet da GVT são exemplos de ADSL.

O ADSL consiste em um sistema de transferência de dados de computador (Internet) usando a estrutura física da linha telefônica (fios, cabos, armários, caixas de distribuição, centrais etc.), usando uma frequência diferente da frequência usada pela linha telefônica (circuito telefônico).

Ou seja, apesar de usar o mesmo fio (ou melhor, par de fios) que a linha telefônica, esse sistema não deixa o telefone ocupado e, por isso, não é tarifado segundo a linha telefônica (pulsos). Esse sistema é usado por várias horas por dia, e o assinante (usuário doméstico ou corporativo) paga apenas uma mensalidade fixa. Eu chamaria esse sistema de “Internet Rodízio”, ao contrário dos dois anteriores que são “Internet self-service no peso”.

O equipamento utilizado por esse sistema é o modem ADSL, que é normalmente fornecido pelas próprias provedoras do serviço (empresas telefônicas) e é um periférico normalmente externo, como visto nas figuras a seguir:



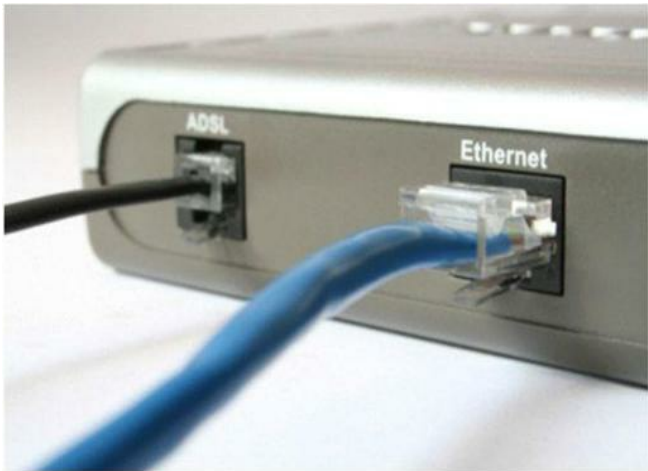
*Figura 9.3 – Sistema ADSL de acesso à Internet.*



*Figura 9.4 – Modem ADSL (Thomson®).*

O modem ADSL pode ser ligado ao gabinete do computador por vários meios, incluindo um cabo USB (esse equipamento pode ser ligado a uma porta USB livre do seu gabinete), mas a forma mais usada e recomendada pelas empresas de telefonia é a conexão através de um cabo UTP com conectores RJ-45.

Se você recorda, esse conector (RJ-45) será encaixado, no computador, na placa de rede do micro, razão por que tantas empresas de telefonia que oferecem o serviço de ADSL exigem uma placa de rede no micro onde o sistema será instalado. Mas se o concurso for bem claro na pergunta: “É necessária uma placa de rede para conectar um microcomputador à Internet por meio de sistemas de banda larga como o ADSL?”, a resposta tem de ser “ERRADO” (pois podemos ligar o modem ADSL pela porta USB).



*Figura 9.5 – Modem ADSL ligado pelo cabo de rede (UTP) com RJ-45.*

A tecnologia ADSL permite velocidades de 64 Kbps a 8 Mbps, em média. Na maioria dos estados do Brasil, são comercializadas as velocidades de 1 Mbps a 8 Mbps.

Um detalhe interessante é que o ADSL é assimétrico (daí seu nome). Isso significa que nessa tecnologia, a velocidade de download (recebendo dados, ou como chamamos: *downstream*) é

diferente da velocidade (taxa de transferência) enviando dados (upload, ou *upstream*). A velocidade de download é sempre maior.

Para se ter uma ideia, quando se contrata o serviço de ADSL de 1 Mbps, essa é a taxa de download, pois a taxa de upload é de cerca de 384 Kbps.

Note que na Figura 9.3 não há menção do provedor de acesso. Isso se dá porque, no sistema de ADSL, quem fornece todos os requisitos para a conexão do usuário (roteadores, endereço IP, DNS etc.) é a empresa telefônica. Ou seja, no sistema de ADSL, a própria empresa telefônica é o provedor de serviços do usuário (embora, no Brasil, ainda sejamos obrigados, por resolução da ANATEL, a pagar uma mensalidade a um provedor – como o Terra, ou UOL –, mesmo que este não nos forneça absolutamente nada!)

O sistema de ADSL usa, normalmente, um protocolo de enlace conhecido como PPPoE (Protocolo Ponto a Ponto sobre Ethernet), que “emula” uma ligação telefônica usando uma rede local. (Os modems ADSL podem ser ligados numa LAN, lembra-se?) Portanto, o protocolo PPPoE é usado para fingir uma ligação telefônica usando a estrutura da LAN, o que é necessário ao ADSL.

#### **9.4.2.1. ADSL 2/ADSL 2+**

Esses são os nomes como são conhecidos os padrões ADSL mais recentemente aprovados pela ANATEL para uso no Brasil (esses serviços já são fornecidos por quase todas as operadoras telefônicas do país).

Através do ADSL2, é possível atingir velocidades de 24 Mbps para download (downstream) e até 1 Mbps para upload (upstream). Como é possível ver, o aumento da velocidade beneficiou o processo de download.

Já é possível adquirir esse serviço na maioria das capitais do país, com velocidades de 10, 15 e até 20 Mbps.

#### **9.4.3. Internet a cabo**

É um sistema de comunicação fornecido por empresas de TV por assinatura (TV a cabo, mais precisamente, não as TVs via satélite). Nesse sistema, os dados de Internet são recebidos pelo mesmo cabo por onde trafegam os dados de TV (de forma similar ao ADSL). Ou seja, o ADSL está para a linha telefônica assim como a Internet a cabo está para as TVs a cabo.

As velocidades do sistema de Internet a cabo são semelhantes às do ADSL (já na nova versão, ADSL2+), incluindo a assimetria (termo que designa o fato de as velocidades de downstream e upstream serem diferentes). Nessa rede usamos um equipamento chamado cable modem (modem a cabo), que é similar ao modem ADSL, mas só funciona nesse sistema.

No caso das conexões por Internet a cabo, o cable modem é conectado ao receptor de TV, que recebe o cabo coaxial que chega da rua (da empresa de TV).

O serviço mais conhecido no país é o Vértua, fornecido pela empresa NET. A NET comercializa acessos residenciais de até 100 Mbps.

#### **9.4.4. Internet através de uma rede local**

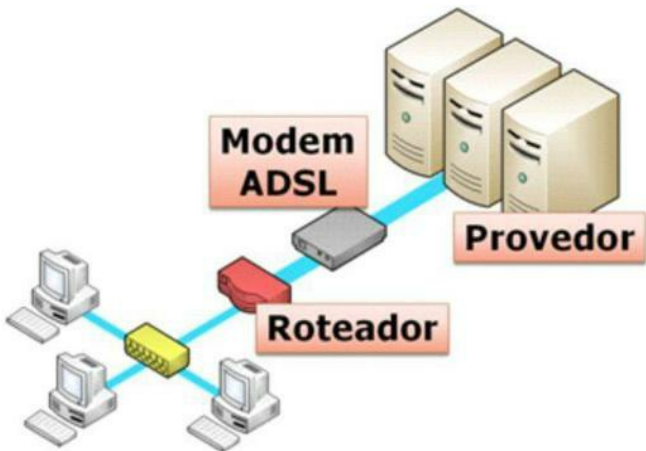
É uma forma muito comum de conexão em empresas e condomínios atualmente. Consiste,

simplesmente, em interligar os micros através de uma LAN Ethernet normal (com hubs ou switches, cabos UTP e placas de rede) ou de uma LAN Wi-Fi (com pontos de acesso, placas de rede Wi-Fi etc.), ligando-as a um roteador que será conectado à estrutura de Internet (que poderá ser qualquer uma das que vimos até agora, como ADSL ou CABO).

Os micros dos usuários precisarão, apenas, de uma placa de rede para se conectarem à LAN e, através dela, terão acesso à Internet por intermédio do roteador. Um roteador é, para que você se lembre, um dispositivo que interliga redes distintas. (Portanto, é ele que tem de ser usado para ligar a sua rede – na sua casa – à rede do seu provedor de acesso.)

Em suma, se você tem mais de um micro para ser ligado simultaneamente a uma única conexão com a Internet, recomenda-se usar um roteador, já que “mais de um micro” já forma uma rede, não é mesmo? Além disso, quando se contrata um serviço de Internet de banda larga, normalmente se ganha do provedor apenas um único endereço IP, ou seja, o suficiente para apenas um único equipamento ser conectado.

O roteador é esse equipamento. Através do protocolo NAT, um roteador recebe o endereço IP externo, oriundo do provedor, e “esconde” os endereços IP internos da rede (endereços que serão usados pelos micros daquela rede). Sendo assim, a Internet inteira vê, na sua casa, apenas um endereço IP (pertencente ao roteador) e não verá os demais endereços IP (dos micros internos).



Para aumentar o desempenho geral da rede no tráfego de dados para a Internet, pode-se usar, além do roteador, um servidor proxy (veremos depois). Nas conexões residenciais (em condomínios, por exemplo), faz-se uso, normalmente, apenas dos roteadores para ligarem a rede inteira ao modem ADSL (ou modem a cabo, dependendo da tecnologia usada para a conexão).

Com certeza, um dos pontos negativos em relação a esse sistema de conexão é que, quanto mais usuários estiverem conectados à Internet simultaneamente na LAN, mais lenta será a comunicação sentida por cada usuário. (Afinal, o fluxo de dados será dividido entre os usuários, se estes estiverem conectados e transferindo dados simultaneamente.)

E lembre-se, também: hoje em dia há equipamentos que mesclam as funções de vários outros. Por exemplo, há vários modems ADSL que também são roteadores e, além disso, são também pontos de acesso e switches. Portanto, os três equipamentos mostrados na figura anterior (hub, não citado, roteador e modem ADSL) poderiam, na verdade, ser apenas um.

#### **9.4.5. Internet através da tomada elétrica**

Um sistema realmente muito novo, conhecido pela sigla **PLC** (Power Line Communication – Comunicação pela Linha de Energia). Os sinais de Internet serão transferidos pela linha de energia normal (aquela que vem da empresa elétrica do seu estado) em uma frequência diferente da frequência da energia que recebemos (60 Hz).

No caso desse sistema de conexão com a Internet, o provedor será a própria empresa de energia elétrica, que ainda deverá passar por algumas atualizações de equipamentos e pessoal para poder fornecer esse novo serviço.

Com esse sistema, os computadores serão ligados a um “modem” ligado à própria tomada onde o computador é ligado à eletricidade, ou seja, não há a necessidade de nenhum conector extra (como RJ-11, RJ-45 ou semelhantes). Estima-se que esse serviço apresente velocidades semelhantes aos serviços de ADSL e Internet a cabo (e preços também).

O sistema, apesar de ser testado há muito tempo, não tem previsões concretas de efetivamente passar a ser utilizado por causa de problemas estruturais nos transformadores dos postes (as ondas que transferem os sinais têm de passar por eles).

#### **9.4.6. Internet via satélite**

Algumas localidades do país já experimentam a conexão direta usuário-provedor via satélite. Esse sistema é necessário quando as empresas telefônicas não podem fornecer o ADSL (devido à distância entre central telefônica e a casa do usuário), nem as empresas de TV por assinatura podem fornecer seu serviço de conexão.

Nesse sistema, o computador do usuário é ligado a uma antena (parecida com as antenas das TVs via satélite), que se conecta diretamente com o satélite da empresa provedora. Esse sistema é normalmente muito caro e sua velocidade é muito boa.

O principal problema do sistema via satélite é a sua instabilidade. Assim como as TVs via satélite, quando o céu fica escuro (nuvens de chuva carregadas), a conexão com o satélite sofre com velocidades baixas e problemas de comunicação.

#### 9.4.7. Internet a rádio

Sistema usado em cidades do interior e bairros mais afastados nas regiões metropolitanas das capitais do país. Através desse sistema, o provedor de acesso se conecta à casa do usuário por meio de antenas.

A conexão não é via satélite, é de antena (em um bairro) para antena (em um outro bairro). O alcance das antenas pode ser de alguns quilômetros, e as velocidades são variadas, mas já considerado banda larga.

#### 9.4.8. Internet via rede celular

Em um momento de convergência de tecnologias de telecomunicações, é comum encontrar vários usuários utilizando seus próprios telefones celulares como “clientes” de e-mail e de páginas, de modo que possam trocar mensagens de correio e acessar a Web sem precisar de computadores oficiais.

Mesmo com um computador convencional (como um notebook) e não um telefone celular, é possível ter acesso à rede de telefonia celular para a transmissão de dados de Internet por meio dessa tecnologia.

As redes de telefonia celular evoluíram bastante no que concerne à transmissão de dados desde os primórdios desse serviço. Inicialmente havia o WAP (Wireless Application Protocol – Protocolo de Aplicações Sem-Fio), que permitia toscos 14,4 Kbps de transferência nas antigas redes de celulares TDMA (os primeiros, de “banda A”).

Depois, as operadoras de telefonia celular escolheram entre duas opções para a então “nova geração”: GSM (o sistema dos “chips”, muito comum na Europa e aqui no Brasil é fornecido pela TIM, Claro e Oi – e agora a Vivo também) e o CDMA (muito comum na Ásia, ideal para transmissão de dados – aqui no Brasil, só a Vivo opera com essa tecnologia).

Cada tecnologia de telefonia celular (GSM ou CDMA) tem suas próprias tecnologias de transmissão de dados de computador (para acesso à Internet, primordialmente). Vamos analisar apenas as tecnologias do padrão GSM, que é o mais utilizado no Brasil:

- **GPRS (General Packet Radio Service – Serviço Geral de Pacotes via Rádio):** é a tecnologia considerada 2,5 G (“segunda geração e meio”) das comunicações de dados nas redes GSM. O GPRS permite conexões com velocidades de 115 Kbps (na prática, 56 Kbps).
- **EDGE (Enhanced Data GSM Environment – Ambiente Melhorado de Dados GSM):** é uma melhoria do GPRS, que permite transmissão de dados a velocidades de 384 Kbps.
- **HSDPA (High-Speed Downlink Packet Access – Acesso de Alta Velocidade para Baixar Pacotes):** é a terceira geração (3G) para o sistema GSM, fornecido pelas operadoras de telefonia celular atuais no Brasil (a maioria delas). Essa tecnologia promete até 10 Mbps de taxa de transferência (7,2 Mbps efetivamente aqui no Brasil) – embora sejam vendidos planos de **1 Mbps** de velocidade normalmente.

Já são muito comuns os “modems 3G” nas operadoras de telefonia GSM: esses dispositivos nem deveriam ser chamados de modem, porque, como vimos no capítulo de hardware, um modem é um equipamento que traduz sinais digitais em analógicos e vice-versa. Os “modems 3G” usam uma rede totalmente digital; portanto, nenhum tipo de tradução é feito.

O nome correto, do ponto de vista conceitual, seria Placa 3G ou **Interface 3G**. E, exagerando no preciosismo, o nome correto seria Interface HSDPA, pois esse é o nome da tecnologia 3G das operadoras do país.



Figura 9.7 – “Modem” 3G ligado à porta USB (interface HSDPA).

Recentemente, em dezembro de 2012, a operadora Claro anunciou o início de suas operações na 4G (4ª geração) em algumas cidades do país (Recife foi a primeira capital a utilizar esta tecnologia).

A tecnologia de 4G da claro é a **LTE (Long Term Evolution – algo como “Evolução de Longo Prazo”)**, que promete conexões com celulares e modems próprios a até 100 Mbps.

### 9.5. Como funciona a Internet

Depois de nos conectarmos à Internet, tendo escolhido o provedor e o sistema que mais nos agrada, o que fazer? Para que a Internet serve realmente? Bom, em grande parte dos casos dos usuários domésticos, o acesso à Internet tem um motivo simples: obtenção de informações.

Quando nossos computadores estão on-line (conectados à Internet), podemos obter diversas informações, bem como fornecer outro tanto delas. É claro que, na quase totalidade dos casos, nossos micros estão conectados para receber informações e não fornecê-las.



Mas, se nossos micros querem receber informações, há quem tenha de fornecê-las, não é? É justamente para isso que servem os computadores intitulados servidores, como já foi visto anteriormente.

## 9.6. Modelo Cliente/Servidor

A “relação” que os computadores mantêm entre si na Internet segue um modelo bem definido, conhecido como modelo cliente/servidor. Nesse “paradigma”, os clientes são os nossos micros, que sempre estão “requisitando” algo e os servidores são os computadores na Internet com a responsabilidade de fornecer esse “algo”.

Vê-se, portanto, que o modelo cliente/servidor é centralizado e é hierárquico. Pois todas as informações estão presentes no servidor que as fornece a quem pede, simplesmente. É como uma sala de aula: os clientes são os alunos, detentores do direito de requisitar informações ao indivíduo que as possui e tem a obrigação de fornecê-las (o professor – no caso, servidor).

A relação entre cliente e servidor é muito simples:

1. O servidor está de prontidão, aguardando que algum cliente solicite algo.
2. O cliente pede uma informação ao servidor apropriado (localiza-o pelo endereço).
3. O servidor então, respondendo à requisição feita, fornece as informações pedidas ao cliente que as solicitou.

Na verdade, os termos “servidor” e “cliente” não definem máquinas no sentido real da palavra, ou seja: o termo servidor não é atribuído a um computador, mas a uma aplicação (programa, software) que tem a função de fornecer informações de diversos tipos. Da mesma forma, um cliente é um programa, não um micro, que foi criado para estabelecer uma conexão com um servidor e dele obter informações.

Como um exemplo: os programas que usamos na Internet são clientes. O Internet Explorer (mostrado na figura a seguir), que é o aplicativo usado para navegar nas páginas, é considerado um programa *cliente web* ou *cliente WWW*, pois se conecta aos servidores Web para obter páginas. O Outlook Express também é uma aplicação cliente, só que cliente de correio ou cliente de e-mail, pois se conecta aos servidores de e-mail para enviar e receber mensagens de correio eletrônico.



Figura 9.8 – Programa navegador (um cliente Web) Internet Explorer 9.

Há diversos servidores envolvidos com os diversos processos que a Internet oferece aos clientes (também há diversos clientes, um para cada tipo de servidor). Há servidores para enviar e-mail, outros para receber e-mails, há alguns servidores para fornecer arquivos e assim por diante. Basicamente para cada serviço oferecido na Internet, há um programa servidor apropriado (e, é claro, um programa cliente apropriado). Veja alguns:

- **Servidor de Páginas:** servidor responsável por armazenar as páginas da WWW (páginas Web) que a empresa mantém em seu site para que os diversos navegadores da Internet consigam visualizar. Esse servidor é chamado, também, de **servidor Web**. Para usar os serviços de um servidor Web, é necessário possuir um programa cliente Web, também

conhecido como **browser** (navegador), como o Internet Explorer (mostrado na Figura 9.8).

- **Servidor de Entrada de E-mails:** também conhecido como servidor de recebimento, ou **servidor POP**, é o servidor, em uma empresa, responsável por armazenar todas as mensagens de correio eletrônico que chegaram à empresa, destinadas aos diversos usuários do serviço (funcionários da empresa). No caso dos provedores de acesso, o servidor POP armazena as mensagens que os clientes do provedor (nós, usuários caseiros) receberam.

- **Servidor de Saída de E-mails:** também conhecido como servidor de envio ou **servidor SMTP**. Ele é responsável por enviar, para a Internet, todas as mensagens de e-mail oriundas dos usuários da empresa. Portanto, cada vez que você, internauta, envia um e-mail pelo seu provedor (Hotmail, Terra, UOL, IG, qualquer um desses), está usando os serviços do servidor de saída do seu provedor.

O programa que usa os serviços dos servidores de saída e entrada de e-mails é chamado **cliente de correio eletrônico**, ou **cliente de e-mail**. Um exemplo interessante é o programa Mozilla Thunderbird, que pode ser baixado gratuitamente da internet e que permite que o usuário se conecte a servidores de saída para enviar seus e-mails e também a servidores de entrada, de modo que o usuário possa receber seus e-mails em seu computador:

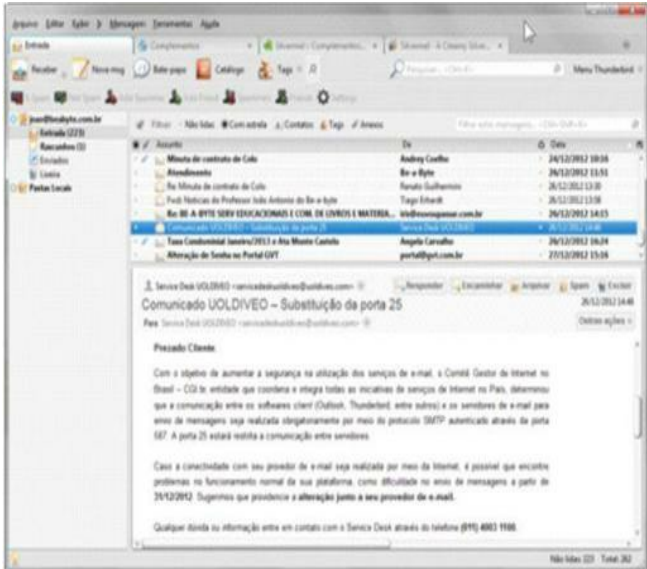


Figura 9.9 – Programa cliente de correio Mozilla Thunderbird versão 17.

Outros servidores comuns:

- **Servidor FTP:** responsável por fornecer uma pasta (diretório) para que os usuários possam acessar e armazenar seus arquivos ou arquivos disponibilizados por outrem. É muito comum a denominação **Servidor de Arquivos** para esse tipo de servidor. Para fazer uso do que esse servidor oferece, o usuário deve possuir um programa cliente de FTP.
- **Servidor Proxy:** é um servidor que realiza a função de “mediar” as comunicações da rede da empresa com a Internet. Quando um dos clientes da rede tenta acessar uma página na Internet, por exemplo, este acessa, na verdade, o servidor proxy, que, por sua vez, acessa a Internet em vez do cliente em si.

Então, colocar um servidor proxy na “porta” de uma rede pode fazer com que o fluxo de informações que entram na rede diminua. Isso porque quando um usuário A pedir uma página,

não pedirá ao site diretamente, mas ao proxy. O proxy, por sua vez, vai à Internet e pede a página, trazendo-a para si (armazenando-a em seu disco), para depois entregá-la ao usuário A!

Se, dentro de alguns instantes, um usuário B, em outro micro da rede, pedir aquela página novamente, ela será pedida ao proxy, que já a possui, enviando-a diretamente ao usuário B sem precisar trazer o grosso do conteúdo da Internet.

Mas um proxy também é usado para auditoria! Sim! Tudo o que passa para a Internet passa pelo proxy. É fácil, portanto, ter acesso ao que os usuários andam acessando simplesmente analisando o histórico (log) do proxy.

Existem vários outros serviços na Internet que são oferecidos, claro, por servidores. Mas esses são os mais importantes. Lembre-se de que para cada serviço oferecido, há um servidor em funcionamento e, também, um programa cliente que possa conversar com ele.

Note também: a nomenclatura de alguns servidores pode ser dada de acordo com o protocolo de aplicação usado para acessar aquele servidor (como em servidor POP, SMTP e FTP). Isso se dá, pelo fato de que cada serviço (tarefa) é acessado por meio de um protocolo específico.

Uma última nota: como cada servidor (aplicação) tem de ser instalado em um computador que, doravante, se chamará servidor, então essa aplicação estará associada a um endereço IP (no caso, o endereço IP do computador). Portanto, lembre-se: um servidor está, necessariamente, associado a um endereço IP, e para se ter acesso aos serviços que este oferece, deve-se conhecer seu endereço IP!

## 9.7. Domínios – nomes amigáveis

Já imaginou ter de decorar os endereços IP de todos os servidores que detêm páginas que você deseja acessar? Por exemplo: 200.211.34.241 para buscar aquelas resoluções de questões de diversas provas da Esaf para a área fiscal. Ou ainda 200.15.16.132 para enviar cartões virtuais de Natal e aniversário aos amigos internautas. E tem mais uma! Que tal acessar 155.254.233.8 para fazer um DOC do seu banco para a conta do seu irmão?

Horível! Seria realmente complicado ter de acessar os principais serviços da Internet pelos endereços IP dos servidores. Mas todos nós fazemos isso (só que não sabemos). Para facilitar a localização de informações na Internet, foram criados endereços que funcionam como “máscaras” ou “maquiagens” para os endereços IP. Esses endereços são conhecidos como **domínios**. Um exemplo de domínio é ***euvoupassar.com.br***.

Ou seja, os domínios, quaisquer que sejam, são, na verdade, endereços associados aos endereços IP dos servidores da Internet. Isso significa que ***www.euvoupassar.com.br***, na verdade, é um “nome bonito” para 198.7.60.161, assim como ***www.redegir.com.br*** é um nome agradável que está associado ao IP 198.7.60.171.

“E o responsável por isso é o DNS (visto no capítulo de Redes), não é João?”

Perfeito, leitor! Isso mesmo! O DNS (Domain Name Services) presente na pilha TCP/IP, que, para alguns, é um protocolo, para outros é um serviço (a maioria aceita que são os dois), é o componente responsável por estabelecer e manter uma estrutura de nomes amigáveis associados aos endereços IP dos diversos servidores da Internet.

Em toda empresa há um **Servidor DNS**, que contém um registro completo dos endereços IP dos servidores e dos nomes de domínios associados a eles. Um servidor DNS é, em poucas

palavras, como aquelas centrais de informações nos shopping centers... “Ei, onde fica a loja tal?” “É só descer a próxima escada e dobrar à direita, senhora.”

Em outras palavras: quando solicitamos [www.qualquercoisa.com.br](http://www.qualquercoisa.com.br), é o servidor DNS que traduzirá isso para 200.231.45.109 para que o nosso micro (cliente) consiga localizar o alvo (servidor). O servidor DNS também é chamado *Servidor de Nomes*.

### 9.7.1. Hierarquia dos nomes de domínio

Em primeiro lugar, só para lembrar, os nomes usados na Internet são chamados de domínios. Esses domínios são gerenciados e organizados por um sistema que está em toda a Internet chamado DNS. Isso permite concluir que todos os nomes usados na Internet seguem as mesmas regras, contidas no DNS. Dentre as várias regras do DNS, está a hierarquia dos nomes de domínios.

“Hierarquia, João? Quer dizer que os domínios têm nível?”

Sim, leitor! Eles são hierárquicos – ou seja, os domínios têm uma organização baseada em níveis (um nível abaixo do outro). Vamos a eles:

#### 9.7.1.1. Domínio raiz (.)

O primeiro nível de domínios dos nomes da Internet, por mais incrível que pareça, não é visto. Todos os endereços de domínio da Internet são, necessariamente, terminados num . (ponto) invisível.

Portanto, qualquer domínio *naoseigual.com.br* é, na verdade, *naoseigual.com.br.*! Todos os nomes de domínio do mundo, não importando se terminam em com, em br, em pt, em jp etc. são subordinados à raiz dos domínios.

Tem autores que costumam referir-se a esse nível hierárquico como “nível 0 (zero)”.

Há 13 servidores DNS de nível raiz no mundo: a maioria deles localizada nos Estados Unidos. Esses servidores são responsáveis por encontrar os servidores DNS de nível inferior a eles (ou seja, os servidores DNS responsáveis pelos domínios de 1º nível).

#### 9.7.1.2. Domínio de 1º nível (TLD e ccTLD)

Exatamente abaixo desse ponto invisível vem o TLD (Top Level Domain – domínio de nível mais alto). Cada um deles é uma “ramificação” diferente da árvore dos domínios no mundo (a árvore que tem como raiz o ponto).

Os TLD são os domínios que registram tipos, como .com (instituição com fins lucrativos, como *hotmail.com*), .gov (órgãos governamentais americanos, como *nasa.gov*), .edu (instituições de ensino, como *harvard.edu*), .org (organizações sem fins lucrativos, como *greenpeace.org*) e assim por diante. Todos os TLD são os topos de suas próprias ramificações e têm suas próprias regras.

Normalmente, os domínios localizados em TLDs originais (esses que representam tipos de instituições) são associados a instituições norte-americanas (com certa razão, visto que a estrutura dos domínios do DNS foi estabelecida ainda quando a Internet basicamente existia só naquele país).

“João, você só falou nos ‘.com’! E os ‘.com.br’, o que são? Como são classificados?”

Com o passar do tempo, caro leitor, e com a expansão da Internet para além-fronteiras dos Estados Unidos, foram criadas outras ramificações de primeiro nível: uma para cada país. São os ccTLD (Country Code Top Level Domain – domínio de nível mais alto para código de país).

São justamente as duas letrinhas finais que acompanham os domínios registrados em árvores fora dos domínios americanos. Por exemplo, o nosso conhecidíssimo “**.br**”. Claro que cada país terá sua dupla de letrinhas (seu ccTLD).

Os domínios da França terminam em .fr; no Japão são .jp; na Inglaterra terminam em .uk; na Alemanha são terminados em .de; em Portugal são .pt etc.

Cada ramificação de país tem seus próprios servidores DNS de TLD. Por exemplo, o(s) servidor(es) DNS do cc **.br** sabe(m) localizar qualquer domínio que termine com .br. Na verdade, é a função deste(s).

Então, se você digitar um endereço qualquer que termine com .br (por exemplo, **euvoypassar.com.br**), seu computador vai consultar o servidor DNS para encontrar o endereço IP associado àquele domínio que você pediu.

Note que os servidores DNS (servidores de nomes) são sempre consultados sobre os domínios subordinados a eles, daí a razão de ser um sistema hierárquico. Veja o que acontece quando você, leitor, pede um endereço do tipo **www.qualquercoisa.com.br**:

1. Sua solicitação é enviada imediatamente para o servidor DNS que serve a você (ou seja, o servidor DNS no seu provedor de acesso).

2a. Caso este servidor tenha a resposta à sua solicitação, ele envia para você o endereço IP correspondente ao nome **www.qualquercoisa.com.br**; ou

2b. Caso o servidor de nomes do seu provedor não possua a resposta para o seu pedido (ou seja, não saiba qual é o IP associado ao domínio que você solicitou), ele vai repassar a requisição a alguém mais “conhecedor” que ele: um servidor DNS superior (que pode ser diretamente o servidor DNS **geral.br** ou outro subordinado a este).

3. Quando a requisição passar por vários servidores “que não sabem” e chega ao servidor DNS **raiz.br**, este, com certeza, terá a resposta (pois ele tem o registro de todos os domínios terminados em .br). Mas preste atenção, leitor: normalmente o servidor DNS **raiz.br** não aponta diretamente para o IP do servidor de páginas **www.qualquercoisa.com.br**, mas para o servidor DNS do domínio **qualquercoisa.com.br**. E este, por sua vez, será o responsável para apontar para os IPs de todos os servidores dentro daquele domínio.

“Mas, João... Agora deu um nó! **www.qualquercoisa.com.br** e **qualquercoisa.com.br** são coisas diferentes? Pensei que seriam exatamente iguais!”

Que bom que perguntou, leitor...

### 9.7.2. Domínios versus nomes dos servidores

Uma coisa é um domínio (**euvoypassar.com.br**), outra coisa são os servidores (aplicações) que serão registrados como pertencentes àquele domínio (como **www.euvoypassar.com.br**). Então fica assim:

- **euvoypassar.com.br** é um domínio registrado no DNS do Registro.br.
- **www.euvoypassar.com.br** é o nome do servidor de páginas daquele domínio. Ou seja, esse

é o nome do servidor que guarda as páginas do domínio *euvoupassar.com.br*. Esse endereço é, em poucas palavras, o site da instituição.

Temos mais possibilidades. Já que há vários servidores em um domínio, cada qual responsável por prestar um serviço diferente (como enviar e-mail, receber e-mail, fornecer arquivos etc.), então temos a certeza de que poderá haver muitos nomes diferentes.

- ***pop.evoupassar.com.br*** seria um nome comum para o servidor de recebimento (entrada) de e-mails. Esse é o servidor onde as caixas postais dos usuários do domínio *euvoupassar* estão armazenadas.
- ***smtp.evoupassar.com.br*** seria o nome de batismo do servidor de envio (saída) de e-mails.
- ***ftp.evoupassar.com.br*** seria o nome do servidor de arquivos (que fornece espaço para armazenamento de arquivos e transferência destes arquivos pela Internet).

Note que: os prefixos “www”, “pop”, “smtp” e “ftp” não são obrigatórios. Eles são apenas padrão, mas podem ser mudados perfeitamente. O nome dos servidores é escolhido por quem administra aquele domínio. Por exemplo, “*sair.evoupassar.com.br*” poderia ser o nome do servidor de saída de e-mails, assim como “*site.evoupassar.com.br*” poderia ser o nome escolhido para o servidor de páginas. (Normalmente, não se contraria o padrão, mas não há nenhum tipo de obrigatoriedade nele!)

Portanto, não sei se você notou, leitor, mas o domínio (*euvoupassar.com.br*, no nosso exemplo) funciona como um “sobrenome” que identifica uma família, enquanto o nome (prefixo) dos servidores atua como o nome que os identificará de forma única dentro daquela família.

Como são servidores (aplicações) diferentes, normalmente eles são em computadores diferentes e, nesse caso, seriam em endereços IP diferentes. Portanto, cada um deles deveria constar nos registros do servidor DNS daquele domínio.

Então, leitor, lembre-se disto: a rigor, cada domínio registrado no mundo tem seu próprio servidor DNS. Esse servidor DNS aponta para todos os servidores dentro daquele domínio. (Ou seja, o servidor DNS de um determinado domínio conhece os endereços IP de todos os servidores subordinados àquele domínio!)

### 9.7.3. Registro de domínios no Brasil

Aqui no Brasil, quando alguém registra um domínio, o faz junto ao servidor DNS raiz .br, através do órgão responsável, chamado Registro.br (o registro pode ser feito pela própria Internet, na página do órgão em <http://registro.br>).

Pois bem, vamos partir do pressuposto básico: se o registro está sendo feito no órgão competente da raiz .br, você deduz que o domínio terminará em “.br”, não é?

“Sim, João, isso deu para entender!”

Pois bem, a instituição dona do domínio (digamos, o domínio *qualquercoisa.com.br*) vai escolher o tipo (.com, .gov, .org, .edu etc.) – esses “tipos” não são TLDs aqui no Brasil. São apenas o segundo nível dos nossos domínios, visto que o TLD (nível mais alto) aqui do Brasil é o “.br”.

Para registrar certo “tipo” de domínio, pode ser obrigatória a apresentação de certos documentos comprobatórios do tipo da instituição, como ***.gov.br***, por exemplo, que só poderá ser registrado por instituições governamentais.



Até o dia 1<sup>o</sup> de maio de 2008 era obrigatória a apresentação de um CNPJ para registrar domínios “.com.br”. A partir dessa data, qualquer pessoa física, apresentando apenas seu CPF, pode registrar domínios desse tipo.

Então, vamos lá. Quando alguém registra um domínio abaixo da raiz .br, o faz registrando seu nome (qualquercoisa) e seu tipo (.com) e terá, claro, por estar abaixo desta árvore, o final (.br). Então, você registra, obrigatoriamente, um domínio com três níveis (o nome da sua instituição é o terceiro nível).

É claro que em domínios que não possuem ccTLD, ou seja, aqueles que possuem apenas TLD, como .com e .org, o nome da instituição está no segundo nível, não no terceiro (hotmail.com – onde hotmail é o segundo nível), pois o próprio TLD, o primeiro nível, já identifica o tipo do domínio.

Vamos continuar: quando uma instituição registra um domínio, por exemplo, qualquercoisa.com.br, ele informará, ao servidor DNS do Registro.br, o endereço IP do seu próprio servidor DNS.

É fácil entender, caro leitor. Se o servidor DNS .br aponta para todos os domínios que terminam com “.br”, é fácil deduzir que o servidor DNS do euvoupassar.com.br aponta para todos os nomes que terminam com “euvoupassar.com.br” – os servidores que trabalham dentro desse domínio, como “www.euvoupassar.com.br”.

Não é, portanto, necessário informar ao servidor DNS do Brasil (.br) os endereços IP de todos os servidores do meu domínio (www.euvoupassar.com.br, smtp.euvoupassar.com.br, pop.euvoupassar.com.br etc.); basta registrar seus IPs no servidor DNS do domínio e, no Registro.br, registrar apenas o IP do servidor DNS do meu domínio.

Sendo assim, quando uma requisição parte pedindo www.euvoupassar.com.br de algum computador no Brasil, a requisição para descobrir o IP desse servidor passa pelas seguintes etapas:

1. O pedido para descobrir o IP de www.euvoupassar.com.br é feito ao servidor DNS do provedor do usuário que pediu. Caso este tenha a resposta, aponta direto para o servidor em questão (www.euvoupassar.com.br); caso não tenha a resposta, o pedido é enviado a um DNS superior, normalmente o DNS do Registro.br direto.
2. O DNS do Registro.br não sabe quem é www.euvoupassar.com.br (o site em si), mas sabe quem é euvoupassar.com.br (o servidor DNS daquele domínio) e envia a solicitação para ele.
3. Ao chegar ao servidor DNS do euvoupassar.com.br, o pedido é analisado e resolvido, pois o DNS do euvoupassar.com.br tem, registrado em si, o IP do www.euvoupassar.com.br (que é o servidor web que contém as páginas daquele domínio). A resposta é, então, enviada ao usuário que a pediu.

Depois disso, o IP do servidor recém-descoberto é usado para que se possa enviar pacotes até esse servidor. Com os pacotes prontos, a comunicação acontece normalmente entre os envolvidos, e a página será devolvida e lida pelo usuário (o resto da história você já conhece).

#### 9.7.4. URL – endereço único dos recursos na Internet

Todos os recursos presentes na Internet (mais precisamente nos servidores) são localizados por

meio de um endereço único conhecido como **URL** (Uniform Resource Locator – Localizador Uniforme de Recursos). O URL tem um formato bastante fácil de entender, cuja sintaxe padrão é:

**Protocolo://servidor/caminho/alvo**

Esse exemplo não explica muita coisa, mas este aqui sim:

***http://www.cespe.unb.br/concursos/nacionais/pf2012/edital.pdf***

Onde:

***http*** é o protocolo usado para realizar a transferência do arquivo que está sendo pedido.

***www.cespe.unb.br*** é o nome do servidor onde o arquivo desejado está localizado. A nomenclatura host ou site também pode ser usada aqui.

***concursos/nacionais/pf2012*** é o caminho dentro do servidor. Em outras palavras, são as pastas (diretórios) dentro do servidor que abrigam o arquivo a ser trazido. Nesse caso, a pasta ***concursos*** contém a pasta ***nacionais*** que, por sua vez, contém a pasta ***pf2012***.

***edital.pdf*** é o arquivo (recurso) que se deseja buscar da Internet (é o alvo do endereço). No nosso endereço, esse arquivo está localizado dentro da pasta ***pf2012***.

Se o usuário que deseja o arquivo conhece o endereço IP do servidor em vez do seu nome, pode usá-lo perfeitamente nesse caso, deixando o URL da seguinte maneira:

***http://200.249.117.89/concursos/nacionais/pf2012/edital.pdf***

Levando em consideração, é claro, que o servidor ***www.cespe.unb.br*** está localizado no computador de endereço IP ***200.249.117.89***.

Bem, com isso espero que você tenha entendido um pouco de como se processa o cadastro e o registro de domínios e para que eles servem. O nosso próximo assunto trará mais uma luz aos serviços que a Internet oferece aos usuários.

## 9.8. Serviços da Internet

De que adianta gastar dinheiro mensalmente para ter acesso à Internet? O que ela nos oferece de interessante? Eis algumas das respostas: os serviços de que podemos fazer uso quando estamos conectados à “Grande Rede”:

### 9.8.1. Correio eletrônico (e-mail)

O correio eletrônico é um sistema computacional de troca de mensagens, não em tempo real, entre usuários. Através desse sistema, os usuários conseguem trocar mensagens entre si. Essas mensagens podem ser compostas apenas de texto ou ter figuras, sons, arquivos anexos, entre outros componentes.

Cada usuário cadastrado no sistema possui um local (um diretório em algum computador servidor) onde poderá receber e deixar armazenadas as mensagens vindas de outros usuários. Esse local é conhecido como ***caixa postal***.

Cada caixa postal é identificada por um endereço único, conhecido como endereço da caixa postal ou endereço de e-mail.

O endereço de e-mail apresenta um formato simples de entender: ***usuario@dominio***, sendo que

**domínio** é o nome da empresa (mais precisamente o nome do domínio da empresa) onde a caixa postal está armazenada e **usuario** é a identificação da caixa postal em si (ou, se preferir, da pessoa dona da caixa).

Então, no endereço **joao@euvoupassar.com.br**, temos que **euvoupassar.com.br** é o nome do domínio (ou “território”) da empresa Eu Vou Passar (o site do qual sou fundador e coordenador). Essa parte do endereço permite que se localize o computador servidor onde está a caixa postal **joao**.

Portanto, **joao** é o nome da minha caixa postal (meu “pedaço de terra”) dentro do domínio (“território”) **euvoupassar.com.br**.

Sempre que você quiser mandar uma mensagem de correio eletrônico para mim, use esse endereço e a sua mensagem será entregue na caixa postal **joao**, presente no servidor de recebimento de e-mails do Eu Vou Passar.

#### 9.8.1.1. Funcionamento do correio eletrônico

Como o e-mail é um serviço cliente/servidor, então podemos destacar dois componentes principais nesse sistema:

**Servidor de E-mail** (também conhecido como **Servidor de Correio Eletrônico**) é um programa que tem como principal responsabilidade enviar e/ou receber as mensagens de correio eletrônico pela estrutura da Internet. Lembre-se de que são os servidores que fazem o “trabalho sujo”. O envio de mensagens de correio depende, e muito, dos servidores de correio.

É comum, em algumas bibliografias, usar o termo **MTA (Mail Transfer Agent – Agente de Transferência de Correio)** para designar os servidores de correio eletrônico.

Os servidores de correio são programas instalados em computadores normalmente localizados nos provedores de serviços.

**Cliente de E-mail (Cliente de Correio Eletrônico)** são os programas que usamos em nossos computadores e que nos dão acesso aos servidores de e-mail. Através dos clientes, podemos solicitar o envio e o recebimento das nossas mensagens de correio eletrônico. Um cliente não faz nada, absolutamente nada, sem um servidor (afinal, é essa a grande verdade do paradigma cliente/servidor).

O termo **MUA (Mail User Agent – Agente Usuário de Correio)** também pode ser usado para designar os programas clientes de e-mail.

O Mozilla Thunderbird e o Windows Live Mail são exemplos de programas cliente de e-mail.

Uma típica comunicação através do correio eletrônico envolve, normalmente, quatro componentes: um cliente para solicitar o envio de uma mensagem (ou seja, o remetente); um servidor para realizar o envio; um servidor para receber a mensagem e mantê-la armazenada; e, por fim, um cliente para solicitar as mensagens recebidas (ou seja, o destinatário). Um resumo desses personagens pode ser visto a seguir:



*Figura 9.10 – Componentes do serviço de correio.*

O envio de uma mensagem de e-mail segue algumas etapas simples:

1. O usuário remetente, utilizando seu programa cliente, redige a mensagem e clica no botão enviar desse aplicativo. A seguir, um exemplo da janela de edição do Mozilla Thunderbird sendo usado para essa finalidade.

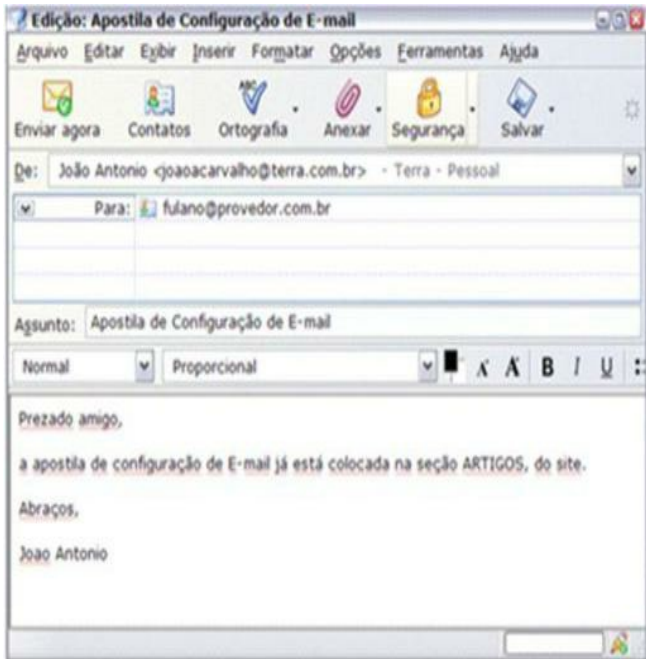


Figura 9.11 – Redigindo a mensagem de e-mail.

2. Quando o usuário solicita o envio da mensagem, seu programa cliente entra em contato com o programa servidor localizado no computador de seu provedor. Esse programa é conhecido como servidor de saída ou servidor de envio, e tem a responsabilidade de enviar as mensagens solicitadas por seus usuários. O protocolo (regra de comunicação) que é usado para esse procedimento é o SMTP (Simple Mail Transfer Protocol – Protocolo de Transferência Simples de Correio), daí o fato de esse servidor ser também chamado de servidor SMTP.

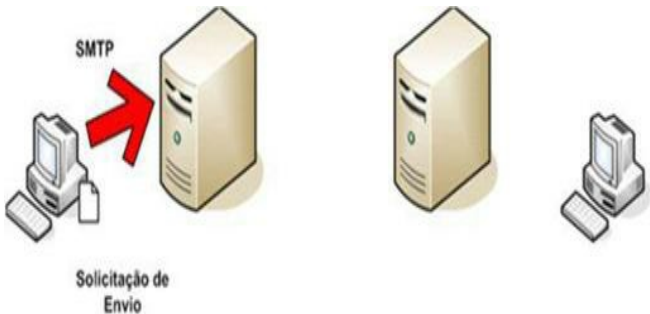


Figura 9.12 – Solicitação do envio – sendo feita ao servidor de saída.

3. Quando a mensagem de e-mail chega ao servidor de saída, este, por sua vez, analisa-a, buscando descobrir para onde ela deve ir. O interessante é que o servidor de saída não se importa com o nome da caixa postal de destino, o servidor apenas analisa o domínio de destino, ou seja, a parte do endereço que está depois do símbolo de @ (arroba).

Portanto, no exemplo anterior, o servidor de saída analisa o endereço que consta como endereço de destino, que, no caso, é *fulano@provedor.com.br*; procura “*provedor.com.br*” em um DNS, para achar o endereço IP associado àquele domínio e, ao encontrá-lo, procede com o envio da mensagem para o domínio *provedor.com.br*. Essa comunicação entre os servidores também se dá por meio do protocolo SMTP.

“Ô, João, por que o servidor de saída não mandou diretamente para a caixa postal de fulano?”

Simples, caro leitor: quando a mensagem chegar lá, no domínio da empresa destinatária, o servidor de lá terá condições de entregar a mensagem a fulano. Acompanhe o restante e veja isso. A seguir, a figura que descreve essa terceira etapa.



*Figura 9.13 – Mensagem enviada entre os servidores.*

4. Quando a mensagem chega ao servidor de destino, conhecido como servidor de entrada ou servidor de recebimento, este se encarrega de analisá-la e, lendo o nome que existe antes do @, que é o nome da caixa postal, armazená-la no local apropriado (o diretório, ou pasta, onde as mensagens daquele usuário devem ser guardadas).

Sim, o servidor de entrada deixa as mensagens armazenadas dentro de si mesmo. As nossas caixas postais não estão em nossos micros clientes, as nossas caixas postais, como foi visto anteriormente; estão localizadas nos servidores de entrada das nossas empresas provedoras.

Um servidor de entrada possui inúmeras caixas postais (pastas) dentro de si. Cada caixa é identificada por um nome único e é por isso que o servidor de entrada consegue determinar para quem a mensagem é enviada. Imagine um servidor de entrada como aqueles armários, em alguns condomínios, que possuem as caixas de correio dos diversos apartamentos.

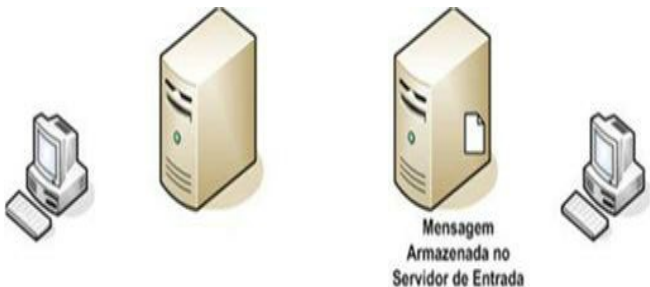


Figura 9.14 – Mensagem armazenada em um servidor de entrada.

5. Depois disso, o usuário destinatário, usando seu programa cliente de e-mail, faz a solicitação ao seu servidor de entrada para que este (o servidor) lhe entregue as mensagens que foram recebidas por ele.

Lembre-se: a solicitação sempre parte do cliente (seja para o envio da mensagem, seja para seu recebimento). Na verdade, o próprio modelo de funcionamento cliente/servidor estipula que a solicitação sempre partirá do cliente.

Ao utilizar um programa cliente de correio, como o Mozilla Thunderbird, as mensagens serão trazidas do servidor de entrada para o micro do destinatário, tirando-as da caixa postal, esvaziando-a (embora se possa configurar o cliente de e-mail para deixar cópias das mensagens na caixa postal no servidor). Essa transferência de mensagens entre o servidor de entrada e o cliente destinatário é realizada pelo protocolo **POP** (Post Office Protocol – Protocolo de Agência de Correio), daí o fato de o servidor de entrada ser conhecido, também, como servidor POP.

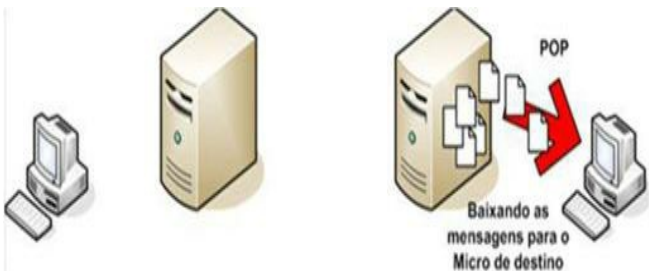


Figura 9.15 – Usuário destinatário recebendo mensagens por POP.

6. Depois de recebidas no computador cliente destinatário, o programa de correio as armazena no disco daquele micro e as mensagens poderão ser lidas, apagadas, respondidas ou encaminhadas a qualquer momento pelo usuário.

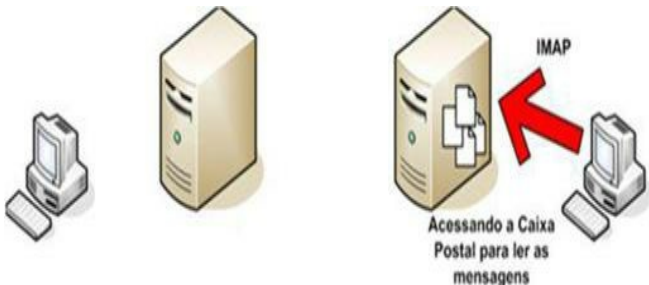
É assim o trajeto de uma mensagem de correio eletrônico. E muita gente pensa que é simplesmente do “remetente para o destinatário”.

Há, porém, uma opção em relação ao recebimento de correio eletrônico por meio do POP: é o uso do protocolo **IMAP** (Internet Mail Access Protocol – Protocolo de Acesso ao Correio da Internet). Esse protocolo é usado pelas pessoas que “pegam” e-mails através das páginas Web de seus provedores (método conhecido como Webmail).

Pois é, quem costuma ter acesso a seus e-mails recebidos por meio das páginas de seus provedores não utiliza POP e, com isso, não recebe as mensagens em seu computador (ou seja, não as traz para seu micro). O IMAP permite que o usuário acesse sua caixa postal diretamente e



leia suas mensagens ainda no servidor de entrada. Qualquer operação, como o apagamento de mensagens, resultará na manipulação de tais recursos diretamente no servidor (ou seja, ao apagar uma mensagem através do webmail, ela será apagada diretamente do servidor, porque, na verdade, ainda estava lá). A figura a seguir mostra o funcionamento do IMAP.



*Figura 9.16 – Usuário destinatário usando IMAP para acessar mensagens.*

Observe apenas uma coisa, caro leitor: através do IMAP, é possível realizar compartilhamento da caixa postal (tecnicamente, também no POP). Ou seja, um usuário cria uma caixa postal qualquer, como “*auditores2013@gmail.com*” e entrega a senha a alguns colegas de estudo, de modo que todos tenham acesso ao conteúdo dessa caixa postal.

Sendo assim, ficam constantemente mandando e-mails para essa “caixa postal geral” de modo que na hora que forem estudar, todos vão pesquisar ali os arquivos que eles vêm guardando há tempos. Essa caixa está funcionando como um “disco virtual”, ou seja, um local onde eles vão armazenar seus dados para compartilhar entre eles.

Outra coisa que é bom entender: por meio do Webmail (que significa, em poucas palavras, “pegar e-mails através de uma página da Web”), os usuários fazem uso, diretamente, do protocolo HTTP, pois têm que acessar uma página (da www) para pegar seus e-mails.

E, como todos já sabemos, acessar páginas da www é coisa para o protocolo HTTP! Mas, vamos à Web, antes de mergulhar mais a fundo nos Webmail!

### **9.8.2. WWW – World Wide Web**

A WWW é um serviço recente na Internet (criado em meados de 1990) que permite que os usuários visualizem documentos diversos na forma de páginas hipermídia. As páginas são arquivos escritos na linguagem HTML, armazenados em diversos servidores espalhados pelo mundo. Esses servidores são chamados servidores Web ou servidores de páginas.

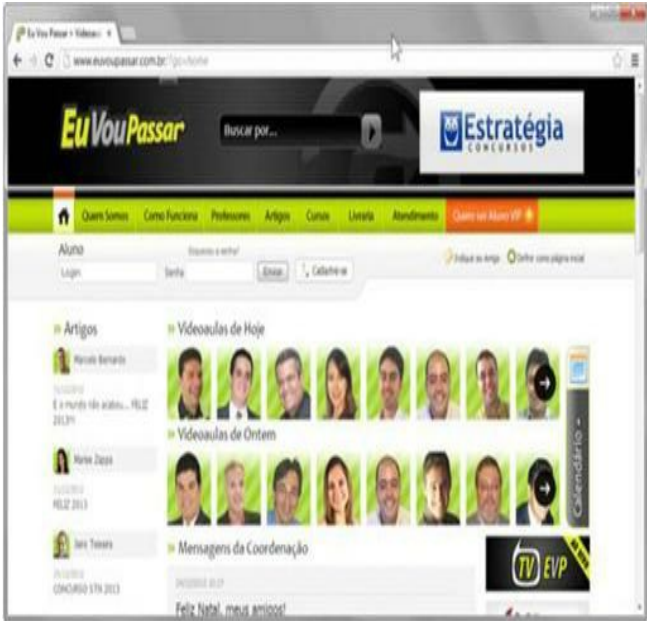


Figura 9.17 – Uma página Web (no site [www.euvoupassar.com.br](http://www.euvoupassar.com.br)).

As páginas são, como já foi visto, armazenadas em servidores, mas a reunião de diversas páginas a respeito de um único assunto ou instituição é chamada **Website** (ou simplesmente **site**). A expressão **sítio da Web** pode ser encontrada também para definir um conjunto de páginas. Pode haver mais de um site no mesmo servidor Web.

A Figura 9.17, por exemplo, está mostrando uma das páginas do site Eu Vou Passar ([www.euvoupassar.com.br](http://www.euvoupassar.com.br)). Portanto, para dirimir quaisquer dúvidas a respeito dos dois conceitos (que, por sinal, muitos confundem), segue um resumo:

- **Página Web:** um documento legível (como uma página de uma revista mesmo). A

definição completa é “documento hipermídia escrito na linguagem HTML”. Na verdade, uma página é apenas um arquivo (sim, um arquivo) que pode ser armazenado em pastas e pode ser copiado de computador para computador.

• **Site da Web:** um local onde são colocadas várias páginas (como se fosse a própria revista). Cada site está associado a um servidor, normalmente. Há, claro, servidores Web com mais de um site sendo fornecidos.

Em poucas palavras, um site pode ser simplesmente uma “pasta” em um computador servidor. Nessa pasta há vários arquivos (as páginas daquele site).

Claro que os sites mais “volumosos”, como Wikipedia, Google, Gmail e outros são formados por muito mais que apenas uma pasta: são milhares de computadores reunidos e espalhados pelo globo!

Portanto, um “site” é uma coisa lógica, que, para existir, pode precisar de muita coisa “física”!

• **A Web (WWW):** é o repositório mundial desses documentos (páginas). Em outras palavras, a Web é o conjunto de todos os sites do mundo (ou seja, a biblioteca).

Lembre-se: páginas são documentos hipermídia. Mas, o que é hipermídia? É um termo que junta dois outros conceitos: hipertexto e multimídia, que serão definidos agora:

**1. Hipertexto:** são textos que apresentam comportamento “ativo”. Ou seja, são textos que não são somente “legíveis” como os textos que encontramos em livros e revistas de papel. Um hipertexto é um texto que permite a existência de hyperlinks (links), aquelas áreas especiais que ficam vinculadas a outras páginas (veremos depois).

**2. Multimídia:** Muitos meios. É um conceito que indica a presença de várias formas de informação (texto, imagem, som, vídeo etc.). Tudo o que for classificado como multimídia está associado à presença de som, vídeo e afins. Como esses tipos de dados podem aparecer em uma página Web, então ela é classificada como multimídia.

**1 + 2. Hipermídia:** é uma definição sobre as páginas Web, visto que elas podem conter vídeo, imagem, som, texto e hyperlinks.

As páginas Web são escritas em uma linguagem conhecida como HTML (Hyper Text Markup Language – Linguagem de Marcação de Hipertexto). Essa linguagem é “universal” para a Web. Veja a seguir um exemplo do código HTML que representa a página mostrada na figura anterior:



O Google Chrome (navegador visto na Figura 9.17) e o Mozilla Firefox também são exemplos de navegadores atuais que podem ser cobrados em prova!

Para visitar um sítio da Web, normalmente digitamos o endereço dele (URL) no campo de endereço, localizado na parte superior da janela do navegador. Depois de acionarmos a tecla ENTER (para estabelecer a conexão com o servidor), nossa requisição é enviada e o servidor nos enviará uma página inicial, que é chamada de Home Page daquele site. Uma **Home Page** é, portanto, a primeira página de um site. É aquela página que aparece imediatamente quando acessamos o URL do site em questão.

Em qualquer página Web, podem aparecer textos, fotos, vídeos e hyperlinks. Dando uma atenção especial a esse último item, um hyperlink (ou simplesmente link) é uma área da página (que pode se apresentar como um fragmento de texto ou uma imagem) que está vinculada a outro documento qualquer. Um hyperlink é, em outras palavras, um “atalho” para outra página, um arquivo, um e-mail ou qualquer outro recurso da Internet.

Reconhece-se rapidamente um hyperlink pela alteração no formato do ponteiro do mouse quando ele está sobre hyperlink (ele vira uma mãozinha apontando para o link).

O protocolo usado na Web para a transferência das páginas é o HTTP (Hypertext Transfer Protocol – Protocolo de Transferência de Hipertexto), que é um protocolo de aplicação da pilha TCP/IP, como já foi visto, e que utiliza a porta 80 sobre o TCP.

“Ei, João, explica de forma fácil: qual a diferença entre HTTP e HTML?”

Simples, leitor: **HTML** é a *linguagem* que é usada para *fazer* páginas. **HTTP** é a linguagem (*protocolo*) usada para *trazer* páginas. Todas as páginas Web são trazidas aos nossos micros por meio de HTTP.

“E todas elas são feitas em HTML, não é?”

Bom... Ai... Nem tanto...

### 9.8.2.1. Páginas estáticas *versus* páginas dinâmicas

Todas as páginas Web são “iguais”? A resposta é, em vários pontos, NÃO! Uma das principais diferenças entre as páginas Web é a forma como elas são interpretadas e apresentadas ao cliente, e as provas da Esaf citam muito e exigem o conhecimento no que classificam como páginas estáticas e dinâmicas. Mas o que essas classificações querem dizer?

Uma página estática é criada em HTML, por alguém de carne e osso (web designer) e é colocada no servidor para ser “mostrada” ao cliente quando este a requisitar. O processo de funcionamento de uma página estática é o seguinte:

1. O Web designer cria a página estática em HTML e a coloca disponível no servidor para os clientes poderem requisitá-la.
2. O cliente solicita a referida página.
3. O servidor recebe a solicitação e fornece a página HTML para o cliente (ou seja, envia uma cópia da página para o computador cliente).
4. O cliente (browser) recebe a cópia da página em HTML e a interpreta, decodificando cada linha do HTML para apresentar a página daquele jeito bonito e agradável com o qual estamos acostumados (textos coloridos, imagens, links etc.).

Viu como é fácil? Em resumo: uma página estática já está presente no servidor quando a

solicitação é feita, porque ela foi criada por alguém que a colocou lá. Se uma página estática for alterada, pode ter certeza de que essa alteração aconteceu porque alguém (o Web designer) a realizou.

Uma página dinâmica é um pouco diferente. As páginas dinâmicas não estão no servidor exatamente como são vistas pelo cliente, quer dizer, o que o cliente vê (código HTML) foi construído no momento em que a requisição foi feita.

Deixe-me explicar melhor: uma página dinâmica é criada com a ajuda de linguagens de programação especiais (chamadas linguagens de servidor, ou linguagens server-sided). Entre as linguagens usadas para esse fim, podemos citar ASP (Active Server Pages, da Microsoft), PHP (gratuita), JSP (Java Server Pages, da Sun Microsystems) e outras mais.

O desenvolvedor da página (Web designer ou Web developer – ih, só nomes para complicar, mas não se preocupe com eles) cria a página na linguagem que achar mais conveniente e a coloca no servidor Web. Quando houver uma requisição daquela página, o servidor irá interpretar o código escrito na linguagem em questão e traduzi-lo para HTML, enviando o resultado para o browser. Fica mais ou menos assim:

1. O desenvolvedor cria a página usando a linguagem ASP, por exemplo, e a coloca no servidor para ficar disponível para os clientes.
2. O cliente, quando quiser, solicita a página em questão.
3. O servidor identifica a solicitação e lê a página. Como ela contém código ASP (linguagem de servidor), a página é interpretada pelo servidor, que transformará todas as instruções escritas em ASP em código HTML (que poderá ser visto pelo navegador).
4. O servidor envia uma cópia da página (devidamente traduzida para HTML) para o cliente.
5. O cliente recebe o código HTML recém-criado e o interpreta, transformando aquele amontoado de instruções na página agradável e intuitiva que vemos o tempo todo.

Portanto, no caso das páginas dinâmicas, o arquivo da página não estava previamente no servidor, mas foi construído no momento da requisição. O que existia anteriormente no servidor era uma espécie de “modelo” da página, escrito em uma linguagem que nosso browser não saberia ler, mas o servidor sabe.

Se não entendeu, pense comigo: como você acha que o seu extrato bancário é apresentado em uma página da Internet? Você acha que um “estagiário” do banco foi contratado para adicionar cada movimentação financeira naquele histórico e colocar as páginas no servidor? Se você acha que sim, como ele faria para fazer isso para todos os clientes do banco?

O que acontece quando você entra no site do seu banco e pede um extrato é o seguinte:

1. O site do banco já contém um “modelo” da página que apresentará o extrato a todos os clientes (sim, você não é exclusivo nesse caso).
2. Quando você entra no site e coloca agência, conta e senha, depois solicita ver o extrato de sua conta, o seu browser (programa cliente) envia a solicitação para que a página do extrato seja mostrada.
3. O servidor entende sua solicitação (e localiza o “modelo” da página solicitada), bem como identifica você como o cliente tal, da conta tal, da agência tal. O “modelo” da página está programado (em ASP, PHP, JSP, ou qualquer outra linguagem de servidor) para acessar o banco de dados da instituição a fim de acessar o histórico de sua conta.

3. Depois de obter o acesso ao seu extrato, o servidor começa a construir o código HTML que apresentará o extrato na sua tela, com as cores predefinidas (créditos em preto e débitos em vermelho, por exemplo). Depois de concluída a construção do código HTML, ele é enviado ao cliente.

4. Seu programa cliente recebe a página recém-construída e a interpreta (claro, é HTML e precisa ser mostrado de forma decente!), mostrando o seu extrato para você e gerando aquele famoso desespero de fim de mês. (Não é um privilégio seu, meu amigo!)

Em suma, páginas estáticas são interpretadas pelo browser (cliente), pois já são escritas em HTML puro, diretamente. Páginas dinâmicas são escritas em linguagens de servidor, por isso são interpretadas pelo servidor e enviadas, já traduzidas para HTML, para o cliente.

#### **9.8.2.2. Cookies**

Já entrou em alguma página Web e foi recebido por ela com “Oi, João!”? (Na verdade, a menos que você também se chame João, tenho certeza de que a resposta é não!) Mas a questão é: como a página Web que estou visitando sabe quem sou eu? Simples: ela “leu” meu crachá!

Quando você acessa uma página qualquer (normalmente de uma loja virtual), deve realizar um processo de cadastro. Esse cadastro consiste em informar ao site alguns dados a seu respeito. Esses dados serão armazenados no servidor daquele site e serão posteriormente consultados quando você acessar novamente.

Mas, para que a página saiba que é você no momento em que o próximo acesso for realizado, ela teve de colocar, no seu computador, um pequeno arquivo de texto com algumas informações básicas a seu respeito (pelo menos o seu número de identificação perante o site). Esse arquivo é chamado cookie.

Um cookie é tecnicamente inofensivo, pois armazena apenas dados relevantes para o site (muita gente pensa que o cookie armazena números de cartão de crédito, conta-corrente etc.), e, se seu computador for invadido e os cookies forem copiados pelo invasor, ele não terá nenhuma informação potencialmente sigilosa sobre você.

Lembre-se: um cookie não é um vírus! Nem sequer pode trazer vírus para o seu computador! Um cookie é meramente um arquivo de texto, colocado no computador do usuário, para identificar aquele usuário em um próximo acesso àquela página.

Nem toda página coloca cookies no cliente, mas, em compensação, há algumas que nem sequer abrem quando não conseguem colocar um cookie corretamente. Sim, é possível um cookie ser rejeitado! No programa navegador, há como configurar o programa para não aceitar nenhum tipo de cookie. Isso é uma ação “paranoica” daqueles viciados em segurança, mas que pode prejudicar a navegação porque certas páginas não aceitam ser vistas em um browser que rejeita os cookies que ela tenta colocar.

#### **9.8.2.3. Webmail**

Apesar de já ter falado rapidamente sobre isso no tópico sobre e-mail, algumas páginas lá atrás, resolvi retomar a explicação sobre Webmail para deixar você sem dúvidas ou qualquer tipo de má interpretação, caro leitor!

Webmail é o nome que nós damos à forma de ter acesso aos e-mails que necessita de uma

página da Web.

Ou seja, se para pegar seus e-mails, você precisa entrar numa página da Web, como do Hotmail, Gmail, Yahoo, entre outros, você está usando Webmail.

Claro que você vai precisar de um browser para isso (um navegador para acessar a página do seu provedor, que dará acesso ao seu e-mail). Não é mesmo? Segue um exemplo:

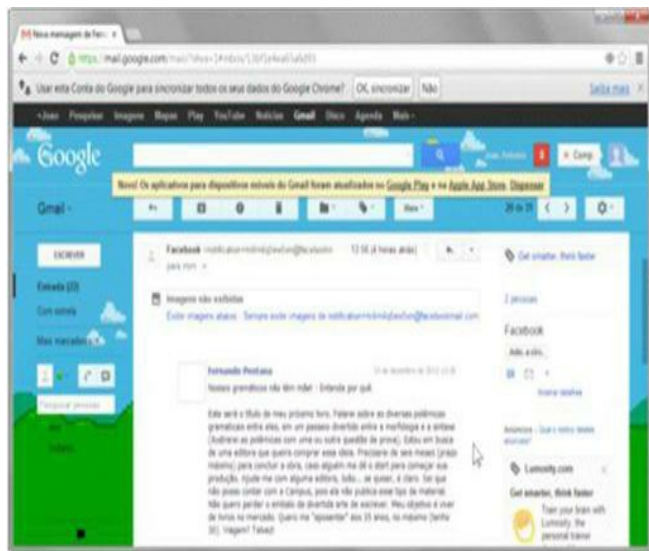


Figura 9.19 – Acessando minha conta no Gmail (usando um browser, claro!).

Pense comigo um pouco, caro leitor: você precisa usar um navegador para acessar uma página da Web para ter acesso ao seu e-mail! Logo, o protocolo que você usa, diretamente, é o HTTP!

Sim! HTTP! Porque a página no seu sistema de webmail é tão página da web quanto qualquer outra página que você acesse! Portanto, se algum elaborador lhe perguntar qual o protocolo usado pelos usuários de Webmail, você já sabe: HTTP!



“Ei, João, e o IMAP, que você mencionou lá atrás?”

O IMAP é usado entre os servidores! Ou seja, entre o servidor de páginas (que lhe atende) e o servidor de e-mails (que efetivamente guarda suas mensagens). Então, você não tem acesso direto ao uso do IMAP, ele é “indiretamente” usado!

#### **9.8.2.4. Redes sociais**

Muito se fala ultimamente em redes sociais, especialmente em concursos. Acho, de todo coração, que é só para assustar os concurseiros, porque o assunto, em si, é muito simples de entender.

Redes sociais são “ligações” entre pessoas e instituições que se relacionam por compartilharem os mesmos interesses ou objetivos. Redes sociais são, em suma, “clubes” que conectam pessoas ou empresas.

A forma mais fácil de ter acesso a redes sociais, hoje, é claro, é por meio da Web. As mais famosas redes sociais do planeta nasceram na Web, por meio da Web (claro que, hoje, há diversos outros meios de acessar tais redes, como aplicativos de celular, videogames e até TVs com acesso à Internet).

As principais redes sociais são, portanto, sites. Vamos a elas:

#### **Facebook**

Quem não conhece o Facebook ([www.facebook.com](http://www.facebook.com))? Rede social criada por Mark Zuckerberg em meados de 2002 e que hoje praticamente “domina” a Internet.

No Facebook, é possível criar grupos de diversos assuntos, estabelecer amizade com várias pessoas (até 5.000 amigos, no máximo, em um perfil), compartilhar fotos, escrever o que quiser, curtir, comentar e compartilhar aquilo que seus amigos postam etc.



Figura 9.20 – Minha Página Inicial no Facebook

Vamos a alguns termos a respeito do Facebook (sei lá, pode ser que caiam em prova):

- **Perfil:** é o cadastro pessoal que você faz no Facebook. Ou seja, para entrar no Facebook, você precisará fazer um perfil (cadastrar-se no site). Cada perfil permite “apenas” 5.000 (5 mil) amigos.
- **Linha do Tempo:** é o nome dado à página inicial de cada perfil. Na Figura 9.20, vemos a minha Linha do Tempo (página inicial do meu perfil).
- **Fanpage:** página especial, usada por pessoas públicas e entidades (empresas). Não é um perfil, é uma página – é uma publicação. Não pode ter amigos, mas pode ter pessoas que “curtam” a página.

Todas as publicações feitas na página são avisadas nos feeds das pessoas que curtiram aquela página. Normalmente, pessoas que chegaram ao limite de seus perfis (5.000 amigos) preferem transformá-los em fanpages.

- **Feed de Notícias:** é a área que nos é apresentada quando entramos no Facebook (quando acessamos o site). Ela contém as mais recentes “atualizações” do que nossos amigos fizeram e do que foi publicado nas páginas que nós curtimos.

Ou seja, não é necessário ficar “garimpando” nas páginas e linhas do tempo dos nossos amigos para ver o que eles fizeram: isso é automaticamente colocado em nosso Feed de Notícias (alimentador de notícias). Ou seja, é “fofoca” novinha que é entregue em domicílio!

facebook  Pesquise pessoas, locais e coisas

 João Antonio

FAVORITOS

- Feed de notícias
- Mensagens 99+
- Eventos 8
- Fotos

ANÚNCIOS

- Gerenciador de anúncios

PÁGINAS

- Eu vou passar.com.br 20+
- João Antonio 20+
- Curta Páginas 20+

GRUPOS

- Eu vou passar.com.br... 20+
- Estudos EvP 7
- Curso básico de Fot... 13
- TRIBUNAIS - Analis... 20+
- Um Dia de Jurista 20+
- ACLMPPF: Alunos d... 20+
- TRT-PR 2012 20+
- ANALISTAS INSS 2012 20+
- Criar Grupo...

Status  Adicionar fotos/vídeo

O que está acontecendo, João?

CLASSIFICAR +

 Paulo Henrique Cremoneze foi marcado na foto de Sara Sanchez. — com Pedrinho Carvalho e outras 24 pessoas.



Usando as belas palavras de nosso poeta Drummond, desejo que em 2013 emerga esse Ano Novo que está desde sempre dentro de você, só esperando uma chance de sair e acontecer. Feliz 2013 a todos nós!

 Bate-papo (Desativado)

Figura 9.21 – Feed de Notícias apresentando publicações recentes de amigos.

Durante muito tempo, o Orkut ([www.orkut.com](http://www.orkut.com)) era a rede social mais usada no Brasil (e o Brasil era o país que mais usava o Orkut no mundo).

Sem muitas “frescuras”, ou seja, com maior velocidade para postagem, criação de grupos e cadastro de amigos, o Orkut era, sem dúvida, a grande “vitrine” das redes sociais no Brasil.

Com a chegada do Facebook, hoje hegemônico, durante algum tempo, via-se o Facebook como “rede social da classe A e B” e o Orkut como sendo “rede social das classes C, D, E...” (Óbvio que esse é o tipo de classificação que não tem nenhum respaldo oficial!).

O que é realidade hoje é: o Orkut está morrendo. Aos poucos, sendo cada vez menos usado... Graças a quem? Ao Facebook! Os usuários têm muito mais recursos no Facebook.. Ele tem tudo que o Orkut tinha, e muito mais!

Ainda há quem use o Orkut... Mas o último que sair, por favor, apague a luz!

### **Twitter – Para os Mais Estressados!**

Twitter ([www.twitter.com](http://www.twitter.com)) é um microblog. Um site onde é possível fazer um cadastro, seguir os cadastros de outras pessoas, e ser seguido por outros usuários.

No Twitter, os usuário escrevem “tweets”, que são pequenas mensagens de texto com até 140 caracteres. Os “seguidores” dos usuários recebem os tweets daqueles que eles seguem em seus feeds.

Ou seja, tudo o que você escreve é lido por quem o segue!

Como usa em sua maioria texto (apesar de permitir, hoje, fotos e vídeos), o Twitter se espalhou rapidamente. As pessoas passaram a tweetar (ou “tuitar”) em qualquer canto (de celulares, tablets, computadores e TVs).

### **LinkedIn – É coisa séria**

Eu, particularmente, sempre achei as redes sociais uma perda de tempo. Nesses sites, a gente escreve o que quer, coloca fotos, compartilha nossa vida com quem a gente não conhece pessoalmente, muitas vezes! Mas há redes sociais para quem é “viciado em trabalho”...

O LinkedIn ([www.linkedin.com](http://www.linkedin.com)) não é uma rede social “com vida social”. É uma rede social profissional, que permite ligar as pessoas por meio de relações profissionais, não pessoais! É quase como um grande banco de currículos de pessoas (empregadas ou não), que se relacionam por alguma ligação de trabalho.

No LinkedIn, só para lembrar, você não tem *amigos*, você tem *conexões*. É possível recomendar suas conexões a outros, é possível escrever resenhas e recomendações profissionais nos perfis das pessoas conectadas a você.

### **9.8.3. Transferência de arquivos – FTP**

Um serviço pouco usado pela maioria dos internautas, a transferência de arquivos através de FTP exige a utilização de um programa cliente FTP e o acesso autorizado a um servidor FTP (servidor de arquivos).

Normalmente, quem utiliza esse serviço são os Web designers para colocar no servidor as páginas que eles alteram constantemente nos sites. Para poder transferir essas páginas, é necessário que eles tenham um login (nome de identificação) e uma senha no servidor.

A transferência de arquivos pela Internet utiliza o protocolo FTP, que faz uso da porta 21 para

os comandos (para cópia, exclusão, renomeação, movimentação etc.) e da porta 20 para a efetiva transferência dos dados. Portanto, duas conexões diferentes são utilizadas para concluir a transferência de um arquivo pelo protocolo FTP. (Sei que repeti isso, mas é só para fixar melhor!)

#### 9.8.4. VPN – Rede Privada Virtual

Este recurso é muito usado em empresas para interligar suas filiais através da estrutura da Internet. Uma VPN (Virtual Private Network – Rede Privada Virtual) é um sistema usado para criar uma rede corporativa (ou seja, pertencente a uma empresa) cujos dados serão transmitidos de forma privada através de uma estrutura de rede pública (adivinha quem? A Internet!).

Mas como é possível transmitir sinais privados através da estrutura física de uma rede pública? Simples: fale em outra língua! É como o exemplo que sempre mostro nas aulas: imagine-se em uma sala de cinema cheia encontrando, no outro extremo da sala, alguém que não vê há anos! Você então resolve contar-lhe um segredo em voz alta (contraditório, não?).

Aí é que está! Como transmitir o segredo no meio da sala de cinema lotada? Grite em outra língua que só você e seu amigo conheçam! A estrutura física que levou os dados (a sala) é pública, mas os dados eram privados e continuarão sendo graças à iniciativa de falar em outro idioma!

Na VPN, os dados que trafegam entre as filias são completamente “estranhos” para o restante da Internet, porque, nesse sistema, se utiliza criptografia (escrita embaralhada dos dados) para garantir seu sigilo. É mais ou menos assim, uma filial envia o dado:

“Conta-corrente: 110.098-8 – Saldo: 12.098,70 – Bloqueado: 23.456,00”

Já imaginou isso trafegando pela Internet assim, desse jeito? Esses dados são muito importantes para caírem em mãos erradas (a Internet está cheia delas!), portanto, são transferidos assim:

II\*&ASK&&777234fj7712QQÇÇ343mmsjuieoosk\*9ksikwy  
eujenbaouJJJOPWUIENRKK34662783JSKNEJROSKEM22939393JJEKSNhs  
jheJSKjJkKkKwWwuioksn!!!!32782jhajh\$\$\$#sjjskjsj

E será traduzido e relido, no destino, assim:

“Conta-corrente: 110.098-8 – Saldo: 12.098,70 – Bloqueado: 23.456,00”

Somente os integrantes da VPN são capazes de entender os pacotes de maneira correta. O restante da Internet é apenas considerado “fora do território” da VPN.

Lembre-se da definição: Uma **VPN** é uma **rede privada que usa a estrutura física de uma rede pública como a Internet**. O funcionamento de uma VPN se baseia em criptografia.

Outra coisa: é comum usarem o termo tunelamento (vem de “túnel”) quando querem se referir à VPN, isso porque os protocolos de uma VPN fornecem o **serviço de tunelamento**, que, no conceito mais simples, é apenas a criação de um “túnel” virtual, permitindo que as mensagens trafegadas por esse sistema não sejam vistas fora desse túnel.

#### 9.8.5. Intranet

Algumas empresas (de grande porte) normalmente criam um ambiente virtual semelhante à Internet: com servidores de páginas para manterem sites, servidores de e-mail para permitir a comunicação via correio eletrônico e até mesmo servidores de arquivos, para FTP. Essa

estrutura visa à obtenção de uma comunicação mais rápida e centralizada entre os funcionários da empresa. Essa estrutura é conhecida como Intranet.

Uma **Intranet** é, no mais simples conceito, **um site interno a uma corporação**. Esse conjunto de páginas é acessível **somente pelos funcionários da empresa** (restrito) e pode ou não ser acessado de fora da estrutura física da rede da empresa.

Usando uma Intranet, os funcionários da empresa podem ter acesso a esse site para encontrar informações pertinentes a eles, podem passar e-mails entre eles e transferir arquivos do interesse da empresa entre seus computadores. Em algumas empresas, a simples existência de um servidor de páginas para manter um site simples (como um “quadro de avisos”) para os funcionários já é tida como uma Intranet.

Em suma, algo imprescindível para a concretização de uma Intranet é a existência de um **servidor de páginas** (servidor Web), porque já é tida como uma Intranet uma estrutura que fornece apenas um **site interno** aos funcionários.

A Intranet utiliza os mesmos protocolos, serviços, portas e aplicativos servidores e clientes que a Internet utiliza. A principal diferença entre as duas é que a Intranet é restrita e, para se ter acesso a ela, é necessária uma autenticação do usuário (login e senha, provando que o usuário é funcionário da empresa).

#### 9.8.6. Extranet

Algumas empresas liberam acesso de parte de sua rede de comunicação interna para pessoas previamente determinadas (funcionários de empresas parceiras, como fornecedores, distribuidores, franquias e até mesmo para clientes finais).

Quando se tem uma parte do sistema de informação da empresa liberado através da Internet para apenas alguns usuários restritamente, tem-se uma extranet. Uma extranet é, em suma, uma parte da rede de uma instituição que pode ter acesso liberado a usuários externos específicos, mediante autenticação por meio de login e senha.

Normalmente utiliza-se esse recurso para permitir o acesso a fornecedores e outros componentes **parceiros daquela empresa na cadeia de negócios**.

#### 9.8.7. VoIP – voz sobre IP

VoIP é o nome que se dá ao sistema que utiliza a Internet, que é uma rede IP (ou seja uma rede baseada em comutação de pacotes por meio de IP), para tráfego de sinais que se assemelham à telefonia convencional (voz em tempo real).

Através do VoIP, dois (ou mais) usuários podem trocar informações em áudio (voz) em tempo real (ou seja, como um bate-papo telefônico normal). Esse termo (VoIP) descreve uma série de tecnologias que permitem que todo o sistema telefônico atual seja um dia transposto para a estrutura da Internet, trafegando por meio de pacotes IP idênticos aos que se utilizam em e-mails e páginas Web.

### 9.9. Principais aplicativos para Internet

#### 9.9.1. Navegadores (Browsers)

São navegadores (ou browsers) os programas que permitem o acesso e a visualização de páginas da Web (www). Estes são, claro, os programas com que mais convivemos no nosso dia a dia:

#### 9.9.1.1. Internet Explorer

Simplesmente o mais conhecido browser do mundo concursário! Pelo fato de acompanhar o Windows, qualquer questão sobre navegadores será, provavelmente, sobre ele!

O IE, como é mais conhecido, está na versão 10 atualmente (acompanha o Windows 8, e, para os que usam Windows 7, pode ser atualizado). Porém, a versão mais provável, ainda, em provas, é o Internet Explorer 9, que originalmente acompanha o Windows 7.



Figura 9.22 – Internet Explorer 9.



### 9.9.1.2. Mozilla Firefox

O Firefox é um navegador criado pela empresa Mozilla. É um software livre e pode ser baixado em [www.mozilla.org](http://www.mozilla.org). É um navegador muito rápido e bastante fácil de utilizar. No início de 2013, quando este livro está sendo editado, a versão mais atual é a 17.



Figura 9.23 – Mozilla Firefox 17.

### 9.9.1.3. Google Chrome

Browser criado pelo Google, o Chrome é o mais novo dos três principais navegadores do mercado. É muito rápido e bastante leve! É o browser que eu mais utilizo!

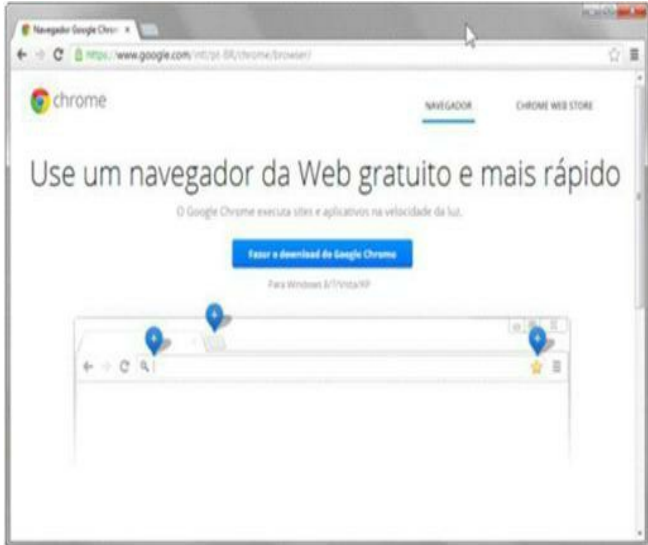


Figura 9.24 – Google Chrome 23 (versão do início de 2013).

Na verdade, quanto ao Firefox e ao Chrome, não há tanta preocupação com versões, porque todo mês, praticamente, se lança uma nova!!! (O Chrome está na versão 23!) Por isso, em provas, provavelmente você os verá serem mencionados, mas sem alusão à versão!

O Internet Explorer não é assim! Como de uma versão para outra leva muito tempo (três anos da versão 9 para a 10, por exemplo), é comum que se façam exigências de versão!

#### 9.9.1.4. Outros navegadores – a galera da “geral”

Ainda podemos citar outros navegadores menos importantes para provas de concursos, mas que podem, eventualmente, ser mencionados:

- **Opera**: navegador gratuito que, apesar de ter uma versão para computadores pessoais, se especializou na versão mobile (para smartphones).
- **Apple Safari**: navegador da empresa Apple, muito comum nos computadores Mac, desta

empresa.

- **Konqueror:** navegador para Linux que acompanha a plataforma KDE.

### 9.9.2. Programas de correio eletrônico

Hoje em dia, a grande maioria dos usuários de Internet não faz uso de programas de correio eletrônico. Eles simplesmente usam Webmail.

Ou seja, até para receber e enviar e-mails, a maioria dos usuários simplesmente utiliza seu navegador para acessar o endereço (URL) do seu provedor de e-mails.

Mas ainda há muito charme e utilidade em se ter um programa de correio eletrônico (um cliente de e-mail) instalado e configurado em seu computador! E essa vantagem se torna muito aparente, especialmente, se você tem um computador que utiliza sempre, e pessoalmente, em sua casa ou trabalho, por exemplo!

Não é muito comum ver perguntas de provas sobre estes programas, afinal, sabe-se que eles não são mais tão utilizados, mas, se vierem, serão provavelmente sobre os aplicativos listados abaixo:

#### 9.9.2.1. Mozilla Thunderbird

Programa livre, da empresa Mozilla. Pode ser baixado do endereço [www.mozilla.com/thunderbird](http://www.mozilla.com/thunderbird). Este é um excelente programa de correio eletrônico, pois oferece uma série de recursos que vão além do simples “enviar/receber”.

Acredito que é muito provável que este seja o programa de correio eletrônico mais cobrado em provas. Por isso, se for instalar um deles, que seja este!

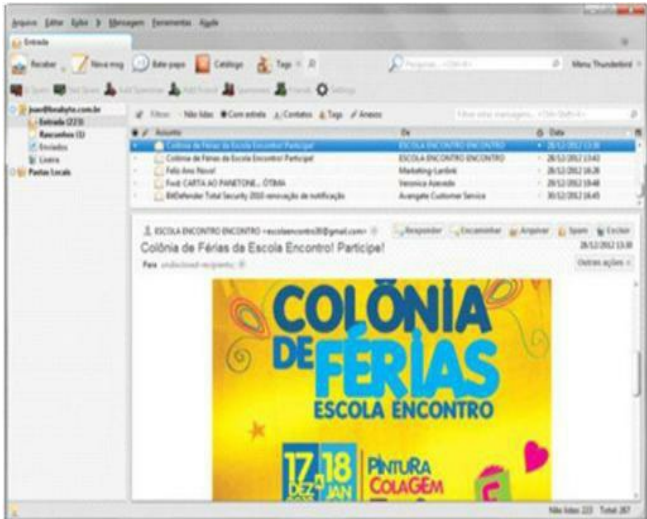


Figura 9.25 – Mozilla Thunderbird 17.

#### 9.9.2.2. Microsoft Live Mail

Antigamente, lá pelos idos do Windows XP, o próprio sistema operacional Windows trazia seu programa de correio eletrônico (na época, o OutlookExpress). Hoje em dia, nem isso!

O programa de correio da Microsoft é baixado da Internet, gratuitamente, num pacote de programas próprio, depois que se instala o Windows. E só baixa se você quiser, caro leitor!

O programa de correio que hoje temos à nossa disposição é o *Windows Live Mail* (é versão “mais bonita e melhorada” do Outlook Express).

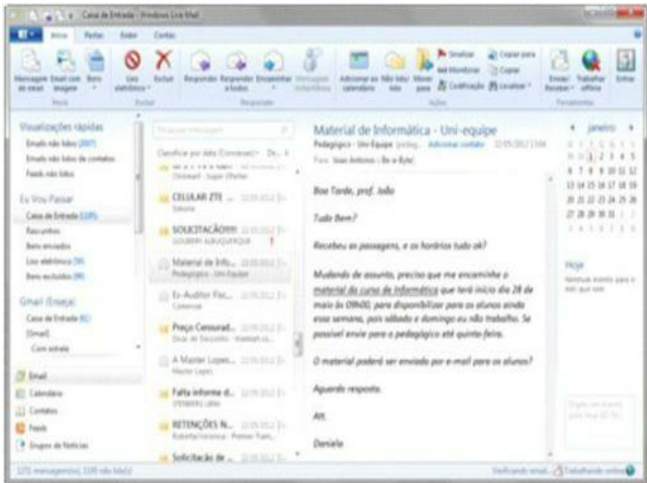
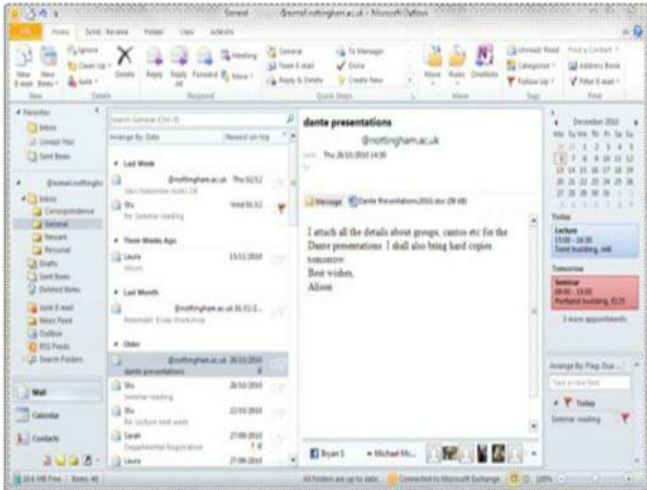


Figura 9.26 – Microsoft Windows Live Mail do Windows 7.

### 9.9.2.3. Microsoft Outlook

Este programa acompanha o conjunto Microsoft Office em algumas versões. É muito mais completo que o Windows Live Mail, por permitir gerenciamento de agenda de compromissos e reuniões, entre outras coisas.



*Figura 9.27 – Microsoft Outlook 2010.*

#### 9.9.2.4. Outros programas de correio

Há outros programas clientes de correio eletrônico possíveis de serem baixados ou comprados. Aqui vão eles:

- **Apple Mail:** programa cliente de correio dos computadores Mac, da Apple (é esse que eu uso, porque meu micro é um Mac).
- **Pegasus Mail:** programa simples e funcional.
- **Evolution:** acompanha o ambiente gráfico Gnome, no Linux.

#### 9.10. Considerações finais

Bem, acho que chegamos ao fim de mais uma etapa. (Por sinal, bastante grande!) Espero que o seu interesse no assunto não fique restrito a esse material (se bem que ele está bastante completo). Existem muitos sites na Internet em que se podem buscar mais informações acerca dos temas.

Qualquer dúvida acerca dos assuntos vistos neste capítulo, não tenha pudor em me mandar um

e-mail (agora que você já sabe como funciona na realidade o sistema). Lembre-se de que joao@euvoupassar.com.br é o endereço da minha caixa postal.

Encontrar-me na Web também é fácil: [www.joaoantonio.com.br](http://www.joaoantonio.com.br) (meu site pessoal) e [www.euvoupassar.com.br](http://www.euvoupassar.com.br). (Simplesmente o mais revolucionário site de auxílio aos estudos para concurso do país – e por que não dizer “do mundo”?)

## 9.11. Q uestões de Internet

1. Observe as seguintes definições:

- I. Aplicativo Web específico para uso interno de uma corporação.
- II. Tráfego Web, comumente no padrão: xxxxxxxx@yyyyyyyyy.com.br.
- III. Tráfego Web, comumente no padrão: http://www.xxxxxxxxxx.com.br.

Correspondem, respectiva e conceitualmente, às definições anteriores:

- a) Intranet; endereço de site da Web e hipermídia;
- b) MS-Word; endereço de site Web e hipertexto;
- c) Internet; hipermídia e endereço de correio eletrônico;
- d) Intranet; endereço de correio eletrônico e endereço de site Web;
- e) MS-Excel; Internet e endereço de correio eletrônico.

2. Considere as propriedades apresentadas a seguir sobre software de correio eletrônico.

- I. Protocolo que permite que mensagens armazenadas em um servidor de correio eletrônico sejam acessadas a partir de qualquer máquina, montando um verdadeiro repositório central.
- II. Protocolo de troca de mensagens entre servidores de correio eletrônico.

Tais propriedades correspondem, respectivamente, aos protocolos:

- a) POP3 e IMAP;
- b) POP3 e SMTP;
- c) POP3 e SNMP;
- d) IMAP e SMTP;
- e) IMAP e POP3.

3. Os protocolos que formam o conjunto TCP/IP são utilizados para atender a uma série de serviços na Internet e em uma Intranet. Com relação aos protocolos que formam o conjunto TCP/IP, é correto afirmar que:

- a) um servidor DNS utiliza o protocolo SMTP para resolver nomes de URLs na Internet e em Intranets;
- b) o protocolo SNMP é utilizado por servidores de e-mail para estabelecer a comunicação com as máquinas clientes no momento do envio de e-mails;
- c) servidores WWW utilizam o protocolo ASP e HTML para estabelecer a comunicação entre clientes e servidores;
- d) o protocolo POP utiliza o UDP para o transporte de mensagens entre estações e servidores;
- e) entre os recursos do IMAP pode-se destacar a recuperação seletiva de partes de mensagens ou mensagens inteiras.

4. Ao configurar-se um aplicativo para receber e-mail, informou-se que o endereço do servidor SMTP da conta de e-mail é smtp.empresa.com.br, equivalente ao endereço IP 123.123.123.123. Após a configuração do aplicativo, utilizando-se o endereço smtp.empresa.com.br, observou-se que este não conseguia enviar e-mail. Em seguida, substituiu-se o endereço smtp.empresa.com.br pelo endereço IP correspondente e verificou-



se que o aplicativo passou a enviar e-mail corretamente. Com relação a essa situação, é correto afirmar que a causa provável do problema está:

- a) no servidor DNS que atende à máquina;
- b) no servidor POP que atende à máquina;
- c) no gateway que atende à rede onde a máquina está localizada;
- d) no roteador que atende à rede onde a máquina está localizada;
- e) na configuração do protocolo SMTP da máquina.

5. Analise as seguintes afirmações relacionadas a conceitos básicos de Internet e Intranet.

I. O POP (Post Office Protocol) é um protocolo que trabalha no ciclo das mensagens eletrônicas. Serve para que os usuários possam enviar facilmente suas mensagens de e-mail para um servidor.

II. O Dial Up é um sistema utilizado pelos browsers para que, quando for solicitado um acesso a um endereço do tipo [www.prova.com.br](http://www.prova.com.br), o computador possa transformar esse nome em um endereço IP válido e realizar a conexão.

III. Um proxy é um servidor que atua como “ponte”. Uma conexão feita através de proxy passa primeiro pelo proxy antes de chegar no seu destino, por exemplo, a Internet. Desse modo, se todos os dados trafegam pelo proxy antes de chegar à Internet, eles podem ser usados em redes empresariais para que os computadores tenham conexão à Internet limitada e controlada.

IV. Protocolos são um conjunto de instruções de como duas ou mais ferramentas se comunicam. O navegador Web e o servidor Web precisam entender um ao outro, por isso os dois se utilizam do HTTP para interpretar as informações que recebem e formular as mensagens que irão mandar.

Indique a opção que contenha todas as afirmações verdadeiras.

- a) I e II.
- b) II e III.
- c) III e IV.
- d) I e III.
- e) II e IV.

6. É muito comum, durante a navegação na Internet, o usuário deparar com sites que se utilizam de cookies, que são:

- a) arquivos que alguns sites criam no seu próprio servidor para armazenar as informações recolhidas sobre a visita do usuário ao site;
- b) arquivos de texto que alguns sites criam no computador do usuário para armazenar as informações recolhidas sobre a sua visita ao site;
- c) vírus especializados em roubar informações pessoais armazenadas na máquina do usuário;
- d) servidores de correio eletrônico que alguns sites utilizam para permitir uma resposta automática a determinadas consultas feitas pelos usuários;
- e) sistemas de segurança utilizados por sites seguros para garantir a privacidade do usuário.

7. Analise as seguintes afirmações relacionadas a conceitos básicos de Internet, Intranet e redes de computadores.

I. Um backbone é a interconexão central de uma rede Internet. Pode ser entendido como uma espinha dorsal de conexões que interliga pontos distribuídos de uma rede, formando uma grande via por onde trafegam informações.

II. Finger é um serviço Internet que permite obter informações sobre usuários de uma máquina.

III. Download é o processo de transferência de uma cópia de um arquivo presente em um computador remoto para outro computador através da rede. O arquivo recebido é gravado em disco no computador local e apagado do computador de origem.

IV. FTP é o protocolo padrão da Internet, usado para transferência de e-mail entre computadores.

Indique a opção que contenha todas as afirmações verdadeiras.

a) I e II.

b) II e III.

c) III e IV.

d) I e III.

e) II e IV.

## Capítulo 10

### Segurança da Informação

#### 10.1. Comentários iniciais

Bem, pessoal, começamos aqui mais um passo no aprendizado da informática para concursos: a segurança da informação, que, apesar de ainda não ser tão cotidiana, passará a ser vista com mais frequência em muitos concursos vindouros.

#### 10.2. Princípios da segurança da informação

Por que segurança? Por que estar preocupado com o meu sistema de computação? Quais os quesitos para classificar meu sistema como sendo seguro? E mais... O que a segurança da informação pode fazer por mim?

Segurança da informação é um termo que descreve técnicas, recursos, componentes e hábitos (sim, hábitos) que permitam que usuários considerem um sistema de informações (um site, um programa de controle de funcionários, um e-mail) confiável.

“Confiável, João? Só isso?”

Você quer mais, caro leitor? Só utilizamos um sistema de informações se realmente entendermos que sua utilização não nos causará problemas, ou seja, se **confiarmos** naquele sistema. A alma do negócio é a confiança! Nem sequer assinamos um contrato sem ter **confiança** na outra parte.

“Mas sempre peço reconhecimento de firma no cartório quando assino um contrato, João!”

Porque você **confia** no cartório, caro leitor! Confia que aquele adesivo e aquele carimbo (e a assinatura do tabelião que você nunca viu) são mais confiáveis que a própria assinatura da pessoa com quem você está fechando o negócio, não é mesmo?

Tudo se baseia na confiança! O objetivo principal da segurança da informação é a confiança que o sistema vai inspirar nos seus usuários! Por isso, aqui vai um termo para você memorizar...

- **Confiabilidade:** descreve a condição em que um sistema de informações presta seus serviços com níveis de eficiência e eficácia aceitáveis. Ou seja, um sistema de informações (um site, por exemplo) irá “desempenhar o papel que foi proposto para si”.

A confiabilidade é, sem sombra de dúvidas, o **objetivo principal** para a existência de recursos e técnicas que visam à segurança das informações.

“Mas como conseguir a confiabilidade, João? Como oferecer aos usuários a certeza de que se pode confiar naquele sistema?”

É fácil (pelo menos teoricamente): basta que o sistema respeite os princípios da segurança da informação. Há quatro princípios básicos, que formam a sigla DICA (ou CIDA, se preferir).

- **Disponibilidade:** é a garantia de que um sistema estará sempre disponível quando necessário (por exemplo, ao acessar um site e ele aparecer, ele estava disponível – se ele não aparecer ou não for possível acessá-lo, o princípio da disponibilidade foi maculado).

Outro exemplo simples da falta de disponibilidade de um sistema acontece sempre no último dia de declarações do Imposto de Renda na Internet, junto ao site da Receita Federal. Tenta deixar para o último dia, para você ver! Acessar o sistema da Receita, só no dia seguinte (com

multa, diga-se de passagem).

- **Integridade:** é a garantia de que uma informação não foi alterada durante seu trajeto do emissor para o receptor ou durante o seu armazenamento. Tendo a garantia de dados íntegros, o receptor pode se assegurar de que a mensagem que ele recebeu tem realmente aquele conteúdo.

Por exemplo, se um e-mail foi alterado antes de chegar ao destino, a integridade foi maculada, mas o receptor não saberia disso até que tomasse a decisão errada influenciada pelo conteúdo fajuto do e-mail.

- **Confidencialidade (Sigilo):** é a garantia de que os dados só serão acessados por pessoas autorizadas, normalmente detentoras de login e senha que lhes concedem esses direitos de acesso.

Também se refere à garantia de que um e-mail, por exemplo, não será lido por outrem a não ser o destinatário devido. (Por exemplo, uma interceptação de um e-mail e a leitura deste por parte de alguém estranho à transação é um atentado à confidencialidade.)

- **Autenticidade:** é a garantia da identidade de uma pessoa (física ou jurídica) ou de um servidor (computador) com quem se estabelece uma transação (de comunicação, como um e-mail, ou comercial, como uma venda on-line).

Essa garantia, normalmente, só é 100% efetiva quando há um terceiro de confiança (uma instituição com esse fim: certificar a identidade de pessoas e máquinas) atestando a autenticidade de quem se pergunta.

Por exemplo, quando você se comunica, pela Internet, com o site do seu banco, você tem completa certeza de que é com o seu banco que você está travando aquela troca de informações?

Quando se puder associar, de forma única e certa, um ato ou documento digital a uma pessoa física (cidadão) ou jurídica, será possível estabelecer regras jurídicas para as transações digitais.

Ainda há dois objetivos secundários, oriundos dos princípios básicos da segurança da informação, são eles:

- **Não Repúdio (irretratabilidade ou irrefutabilidade):** é a garantia de que um agente não consiga negar (falsamente) um ato ou documento de sua autoria. Essa garantia é condição necessária para a validade jurídica de documentos e transações digitais.

Só se pode garantir o não repúdio quando houver **autenticidade** e **integridade** (ou seja, quando for possível determinar quem mandou a mensagem e quando for possível garantir que a mensagem não foi alterada).

Novamente, entramos no mérito de que só haverá tal garantia 100% válida se houver uma instituição que emita essas garantias.

- **Privacidade:** é a condição em que um componente do sistema (usuário) tenha de controlar quem vê as informações sobre si e sob quais circunstâncias.

Deixa-me ser mais claro: ter meu nome e telefone estampados em um outdoor fere a minha privacidade? O que você acha, caro leitor?

“Claro que sim, João! É óbvio!”

E se eu disser que não, leitor? Se fui eu quem determinou que o meu nome e telefone estariam lá, isso não fere minha privacidade. Ou seja, privacidade é a minha capacidade de escolha

(capacidade de controle) sobre as informações.

Se alguém, sem minha autorização, colocou meu nome e telefone no outdoor, aí sim! Isso é falta de privacidade! Porque eu não tive controle sobre o fato de meus dados estarem lá.

Só se consegue efetiva privacidade se o sistema promove a **confidencialidade e autenticidade**, ou seja, se há mecanismos para saber quem é quem (autenticidade) e mecanismos para proibir que alguns “quem” tenham acesso às informações de outros “quem” (confidencialidade).

### 10.3. Ameaças aos sistemas de informação

São componentes que podem prejudicar, de forma temporária ou permanente, o funcionamento de um sistema de informação. As políticas e agentes de segurança têm como principal objetivo evitar que tais componentes tenham sucesso.

- **Defeitos de hardware:** infelizmente, não há como prever tais falhas. O que se pode fazer para evitar que tais problemas sejam muito prejudiciais aos dados do sistema é a realização periódica de cópias de segurança (backups).

Note bem: backups não evitam as falhas! Backups apenas garantem que poderemos recuperar os dados em caso de algum sinistro com nossas máquinas.

- **Hackers:** usuários experientes (conhecedores a fundo) em sistemas de informática. Os indivíduos denominados hackers não são necessariamente ameaças, pois existem os “hackers do bem”.

Apenas são conhecidos pelos seus conhecimentos avançados em informática e, especialmente, redes de comunicação. Alguns poucos indivíduos dessa categoria são capazes de peripécias antológicas, como a invasão de sistemas de segurança da NASA e do Pentágono; portanto, teoricamente, nada os pararia, mas a maioria dos que se intitulam hackers não consegue ultrapassar um firewall bem configurado e um sistema atualizado.

- **Crackers:** usuários experientes que quebram sistemas de segurança (como acesso) ou quebram sistemas de proteção a softwares (senhas e números de série dos programas).

Também não vou aqui exprimir qualquer juízo de valor com relação aos crackers. Eles normalmente estão, sim, “do outro lado da lei”, pois quebram sistemas de proteção de propriedade intelectual (licenças de softwares). Mas o que se pode dizer deles é que conhecem, em profundidade, programação!

- **Programas desatualizados:** os sistemas operacionais e aplicativos apresentam falhas diversas que, com o tempo, “caem na boca do povo”. Quando uma falha é descoberta, os hackers (e os quase-hackers) de plantão saem à procura de sistemas que ainda não foram atualizados e que, por isso, ainda possuem tais falhas.

Manter o Windows atualizado, bem como qualquer outro programa de comunicação com a Internet, é exigência para se ter um sistema menos suscetível a essas falhas.

- **Spam:** envio de mensagens de e-mail em grande número (sem autorização dos destinatários). O spam não é uma ameaça à segurança em si, mas que é chato, é!

Alguns programas, ditos Antispam, tentam diminuir os efeitos dessa prática abusiva, mas muitas vezes sem sucesso. (Os programas filtram quais mensagens devem ser consideradas spam e quais devem ser consideradas mensagens válidas, mas, muitas vezes, não as classificam direito!)

- **Usuários descontentes/leigos:** podem causar problemas com/sem intenção (respectivamente). Quando um usuário não sabe o que está fazendo ou não consegue mensurar a importância de sua senha estar bem guardada, muitos problemas podem acontecer por meio de ataques ao sistema da empresa propiciados pela, digamos, “ingenuidade” do usuário.

A intenção de causar problemas ou de abrir portas para invasores pode ser também fator marcante dentre os problemas que um sistema de informação pode enfrentar.

- **Fraudes/golpes:** técnicas que se utilizam da ingenuidade ou do emocional dos usuários para permitir a obtenção de dados privados de suas vítimas ou para convencê-los a realizarem operações que colocarão em risco a segurança do seu sistema (como baixar arquivos perigosos). Veremos alguns dos tipos de fraudes e golpes mais comuns em um tópico a seguir.
- **Malware:** são programas criados com o intuito de prejudicar usuários e sistemas de informação. Existem vários tipos de malware, e os mais importantes são discutidos no próximo tópico.
- **Ataques:** são atos deliberados de usuários a fim de invadir, destruir ou simplesmente espionar sistemas de informação. Lembre-se de que ataque é um ato deliberado e intencional (doloso). Vamos conhecer alguns dos principais tipos de ataques a sistemas de informação em um tópico seguinte.

### 10.3.1. Malware – programas maliciosos

#### 10.3.1.1. Vírus de computador

Um vírus de computador é um programa (ou parte de um programa) de computador, normalmente com intenções prejudiciais, que *insere cópias de si mesmo em outros programas e/ou arquivos* de um computador, se tornando parte destes.

Um vírus não consegue, em outras palavras, “viver” sozinho! Vírus de computador, assim como os vírus biológicos, precisam de um hospedeiro (que, no caso do computador, é um arquivo qualquer).

Enquanto um vírus encontra-se dentro de um arquivo, ele está em estado de “latência”, simplesmente dormindo. Para que o vírus comece a atuar no computador da vítima, é necessário que o arquivo que o hospeda seja executado (aberto na memória principal daquele computador).

“Quer dizer que se eu simplesmente copiar um arquivo que está infectado com vírus de um pen drive para o meu disco rígido, meu micro não foi infectado?”

Precisamente, leitor! O arquivo que contém o vírus (vamos chamá-lo de “portador do vírus”) pode ser copiado diversas vezes entre vários computadores e nada vai acontecer! Você poderá copiar o arquivo portador do vírus de micros quaisquer para o seu micro. O que não seria muito “legal” para você seria a abertura desse arquivo!

Lembre-se: um vírus é um programa (ou parte de um). Ele é formado por instruções, como todo programa. E, também como todo programa, essas instruções só são executadas (postas em prática) se o vírus for levado à memória principal (RAM). Então, para que o vírus continue seu processo de “replicação e destruição”, ele deverá ser executado no micro.

Depois de executado no seu micro, o vírus, teoricamente, pode fazer qualquer coisa. Qualquer coisa a que foi programado:

Normalmente, a primeira coisa que um vírus faz é se replicar. Ele simplesmente procura outros arquivos nos quais poderá incluir uma cópia sua – é o “instinto de sobrevivência”.

“Mas, João, quando eu desligar o micro o vírus vai parar de ser executado. E quando eu ligar novamente o micro, basta que eu não abra nenhum daqueles arquivos infectados para o vírus não voltar a atuar, não é mesmo?”

Sim, leitor! Mas, como você sabe quais seriam os arquivos infectados para não abri-los? E mais: os vírus são programados para se copiar, especialmente, para arquivos específicos do sistema operacional que sempre são abertos quando o micro é ligado. Isso garante que quando o computador for novamente iniciado, o vírus garanta “seu lugar ao sol” na RAM daquele micro.

Só para você se lembrar:

1. Vírus são programas que se copiam sozinhos para anexarem-se a outros arquivos (os hospedeiros).
2. Vírus precisam dos hospedeiros.
3. Para que o vírus comece a trabalhar (infectar outros arquivos e prejudicar o micro), é necessário que se execute (abra) o arquivo hospedeiro.

#### 10.3.1.2. Worms

Worm (um verme) é um programa capaz de se propagar automaticamente através de várias estruturas de redes (como e-mail, web, bate-papo, compartilhamento de arquivos em redes locais etc.), enviando cópias de si mesmo de computador para computador.

“Então, Worms são vírus que se propagam pelas redes?”

Não, leitor. ***Worms, definitivamente, não são vírus!*** Os vírus conseguem inserir cópias de si mesmos em arquivos, como vimos, tornando-se parte deles. Diferentemente dos vírus, os worms não inserem cópias de si mesmos em outros programas ou arquivos! Os Worms são seus próprios arquivos, ou seja, não precisam de hospedeiros porque possuem corpo próprio.

O objetivo principal dos Worms não é prejudicar ou danificar computadores e/ou arquivos em um sistema, mas, simplesmente, propagar-se. Ou seja, os worms são criados para “passear” pelas redes.

“E isso é prejudicial em algum ponto? Quero dizer... Deixa o cara passear em paz!”

Sim, leitor. É prejudicial, em primeiro lugar, porque gera uma sobrecarga excessiva no tráfego da rede, tornando-a mais lenta. Afinal, copiar-se indiscriminadamente pelas redes de micro em micro vai gerar um tráfego excessivamente grande.

Em segundo lugar, por sua incrível “desenvoltura” em trafegar pelas redes, os Worms podem ser os vetores perfeitos de vírus e outras ameaças (ou seja, podemos “incluir vírus” em um Worm para que este “carregue” aqueles nas costas em suas viagens).

“Ah! Ok! Mais alguma diferença entre eles?”

Sim! Um Worm não necessita ser explicitamente executado pelo usuário-alvo para se propagar. Ele só é executado uma vez (lá no computador que originou a viagem) e ele continuará sozinho, porque sua propagação se dá através da exploração de vulnerabilidades existentes nas redes ou falhas na configuração de softwares instalados nos computadores dessas redes.

#### 10.3.1.3. Cavalos de Troia (Trojan Horses)

Cavalos de Troia são programas, normalmente recebidos de forma aparentemente inofensiva, como por exemplo, uma foto, um jogo, um cartão de aniversário virtual etc., que, além de executar funções de fachada para as quais foi aparentemente projetado, também executa outras operações sem o conhecimento do usuário.

“Quer dizer que ser Cavalo de Troia é ser hipócrita?”

Sim, leitor! Bela definição! A história de Homero conta que uma estátua na forma de cavalo de madeira foi dada aos troianos pelos gregos, como sinal da paz entre eles. Só que esse “presente de grego” demonstrava a hipocrisia dos gregos. Foi um tremendo vacilo!

“Lobo em pele de cordeiro” – isso é a definição de um Trojan!

Um cavalo de Troia normalmente se apresenta na forma de um único arquivo que precisa ser executado pelo usuário para que este sofra as consequências desse ato impensado. Ou seja, o trojan só realizará suas ações (as inofensivas de fachada e as sacanas escondidas) se for executado no micro do usuário vítima.

Depois de executado no computador-alvo, o trojan poderá realizar uma série de ações maliciosas se estiver programado para isso: ele poderá instalar outros programas maliciosos (como vírus, keyloggers e screenloggers); roubar senhas e informações dos usuários, como cookies; modificar ou apagar de arquivos variados; instalar backdoors para que o micro fique sempre vulnerável – com uma porta aberta – para futuras invasões etc.

“Dá até medo perguntar: Cavalos de Troia não são vírus, são?”

Não, leitor! *Cavalos de Troia não são vírus nem worms!* Os trojans não infectam outros arquivos. Um trojan é normalmente um arquivo executável. Esse arquivo tem de ser executado, e ele foi construído daquele jeito – com aquela atuação hipócrita, como você mesmo descreveu. Além disso, os trojans, por definição, não criam cópias de si mesmos autonomamente.

#### 10.3.1.4. Keyloggers e Screenloggers

Um keylogger (algo como “registrador de teclas”) é um programa que armazena todas as informações que um usuário digitou em um micro infectado por esse tipo de programa. Um keylogger é um “presente” muito comum em spywares.

Os teclados virtuais dos sites dos bancos (em que se insere a senha através de cliques do mouse, em vez de teclado) são artifícios criados para evitar a captura de informações por meio de keyloggers, e são bem eficientes contra eles, já que os keyloggers se limitam a capturar os dados inseridos via teclado (ou seja, capturam apenas o que se digita).

Porém, a “galera do mal” também evoluiu! Os “descendentes” dos keyloggers são os screenloggers, ou “registradores de tela”, que armazenam dados quando o usuário clica com o mouse.

Um screenlogger pode armazenar a posição (x,y) do ponteiro do mouse no momento dos cliques. Isso faria o espião, ao entrar no mesmo site em que o usuário estava, deduzir onde os cliques foram dados em cada momento, permitindo que ele faça uma “reconstituição” do trajeto do mouse em cada clique a fim de repetir a sequência de cliques dados no teclado virtual e, com isso, imitar a senha do usuário de quem os cliques foram capturados.

Um screenlogger mais bem elaborado pode, inclusive, capturar a área (uma pequena imagem



da área ao redor do botão) onde o clique foi dado, capturando, assim, em formato de imagem, o próprio local do clique (o botão). Com isso, realmente, lá se vai a segurança dos teclados virtuais, não é mesmo?

“Não, João! Existem sites de banco em que o teclado virtual é aleatório e quando a gente clica nele, todos os caracteres viram \* (asteriscos). Um screenlogger, pelo que você explicou, capturaria apenas asteriscos, não é mesmo?”

Sim, leitor! É mesmo! Muito bem! Os loggers evoluem, mas os teclados virtuais têm de evoluir também. O teclado que você descreveu é perfeito para evitar os screenloggers.

Claro que nunca é demais evoluir os sistemas dos teclados virtuais dos bancos, né? Há casos em que o usuário clica num botão que poderá significar vários caracteres, só o verdadeiro dono da conta-corrente saberá em quais botões clicar!

### **10.3.1.5. Spyware e Adware**

Um spyware é um termo que descreve uma grande gama de programas que monitoram os hábitos de acesso e navegação dos usuários. Um spyware não é necessariamente um programa implantado “ilegalmente”.

Algumas empresas usam spywares nos micros dos seus funcionários (isso deve estar previsto, óbvio, nos termos assinados pelos dois no momento do contrato) para saber exatamente o que os funcionários andam fazendo nos micros da empresa.

Normalmente, quando instalado de forma ilegítima, um spyware está associado a uma série de ações que podem ser realizadas em um micro sem a autorização do usuário, como: o monitoramento e armazenamento dos URLs (endereços digitados nos navegadores) acessados por aquele usuário; a instalação de keyloggers e screenloggers naquela máquina para a captura de tudo que o usuário digita ou clica; monitoramento e captura de informações inseridas em outros programas (não só no navegador), como programas de texto e planilhas; entre outras.

Spywares transformam seu micro em uma “casa do Big Brother”, ou seja, dá pra ver tudo que você faz no seu micro se um spyware for instalado nele.

Já os adware são programas que fazem anúncios de propaganda em seu computador.

Existem casos em que os adware são, até certo ponto, lícitos: quando aparecem dentro de outros programas (gratuitos) a título de patrocínio. É possível encontrar propagandas variadas no MSN Messenger, no Emule, no WinZip e em vários outros softwares que baixamos gratuitamente na Internet.

Mas um adware malicioso mesmo fica à espreita, na memória RAM do seu micro para, quando achar adequado, abrir uma janela de um navegador apontando para a página de um patrocinador. (Pode ser uma página de cassino, venda de remédios, réplicas de relógio de luxo, páginas pornográficas e muito mais – eu já vi de tudo!)

Depois de instalados, os adware podem, sem sua autorização, inserir páginas na sua lista de favoritos, alterar a página inicial do seu navegador (para a página de um de seus patrocinadores), bloquear até mesmo o direito de você, usuário legítimo, alterar aquela página inicial novamente, entre outras coisas tão “agradáveis” quanto estas.

### **10.3.1.6. Backdoor (“Porta dos Fundos”)**

Um programa que, colocado no micro da vítima, permite que o invasor que o colocou possa facilmente “voltar” àquele computador em um momento seguinte.

Um backdoor é uma “brecha” de segurança intencionalmente colocada no micro da vítima para permitir que este tenha sempre uma porta aberta para o invasor poder voltar àquele micro sem precisar utilizar as mesmas técnicas que utilizou na primeira invasão.

Backdoors são trazidos, normalmente, por programas como cavalos de Troia ou enviados por e-mail ou outro meio qualquer normalmente na forma de um único arquivo executável.

Pode-se entender um backdoor como um pequeno programa servidor que habilita um serviço em uma porta específica do seu computador (como a porta 80 está para o HTTP) e permite que o invasor, possuidor do programa cliente correspondente, possa se comunicar com o computador a fim de controlá-lo à distância ou mesmo ler o conteúdo do seu disco.

#### **10.3.1.7. Exploits**

Programas que exploram falhas em sistemas de informação. São programas prontos que os hackers constroem para os que “estão na escolinha de hacker”. Esses programas são criados para utilizar as falhas previamente descobertas nos sistemas.

Quando um hacker (ou cracker) descobre uma falha em algum sistema de informação que possa comprometê-lo, ele normalmente cria um programa (ou parte de um programa) para explorar aquela falha recém-descoberta, a fim de facilitar, para ele e para os que não são tão brilhantes quanto ele, a realização de uma nova invasão.

#### **10.3.1.8. Sniffers (capturadores de quadros)**

São programas que capturam quadros nas comunicações em uma rede local, armazenando tais quadros para que possam ser analisados posteriormente por quem instalou o sniffer.

Um sniffer é completamente efetivo em um único segmento de rede, ou seja, ele será perfeito se entre os computadores envolvidos (o espião e o espionado) houver apenas um hub. Os sniffers se baseiam no recebimento e na não rejeição dos quadros que chegam à placa de rede do computador que está espionando.

Ou seja, um programa sniffer instalado em um computador simplesmente faz a placa de rede atuar em modo promísco (recebendo e processando todos os quadros que chegam a essa placa, mesmo se não forem realmente endereçados a ela).

Se, entre o micro em que está o sniffer e o micro espionado (ou será espionado), estiver um switch, o trabalho do sniffer fica bem mais difícil, porque o switch não vai, naturalmente, usar broadcast para enviar dados a todos os computadores da rede; logo, por ser um filtro natural, o switch enviará os sinais apenas ao micro que deve recebê-los, não enviando os sinais elétricos ao micro espião.

Para conseguir efeitos em redes com switches, alguns programas sniffers prometem enganar o switch adulterando suas tabelas de endereços MAC, fazendo com que os switches apontem para si os quadros que deveriam ser enviados a outrem. Para fazer isso, os sniffers usam técnicas como o MAC spoofing (para alterar o endereço MAC dos quadros que saem do computador que possui o sniffer).

### 10.3.1.9. Port Scanners

Programas usados para varrer um computador para saber quais serviços estão habilitados naquele micro que se deseja invadir. Ao ato de varredura, em si, é chamado de Port Scan.

Usando um port scanner em seu micro, um invasor pode ter certeza de quais serviços estão sendo servidos no micro-alvo e, com isso, pode desenhar a melhor estratégia para a invasão com base nas portas que estão abertas (lembre-se de que cada serviço é prestado através de uma porta específica).

Normalmente, os port scanners não participam da invasão em si, mas são peças-chave no processo de pré-invasão (ou seja, na preparação para o processo de invasão).

### 10.3.2. Fraudes e golpes na Internet

#### 10.3.2.1. Phishing (ou Phishing Scam)

É um golpe muito utilizado para obter dados de usuários desavisados ou fazê-los abrir arquivos com programas maliciosos.

Consiste em enviar aos usuários (normalmente por meio de e-mail – em algum spam) uma mensagem ilegítima que aparenta pertencer a uma instituição conhecida, como um banco, ou um órgão do governo (Receita Federal, INSS e Ministério do Trabalho são apenas alguns dos que eu já recebi).

Nesses e-mails falsos, há normalmente links que apontam para páginas falsas que nos pedem nossos dados (nome, CPF, número da conta e, claro, senhas). Alguns desses links também são usados não para nos levarem a páginas ilegítimas, mas também para que baixemos arquivos (perigosos, claro) para nosso computador.

Uma das formas de evitar ser enganado por esse tipo de técnica é, ao receber um e-mail de qualquer instituição da qual você faz parte (por exemplo, um e-mail do Banco do Brasil com uma oferta tentadora, mesmo que você seja correntista de lá), ***não clicar em nenhum link daquele e-mail.***

Em vez disso, vá ao site do referido banco ou empresa (digitando o URL dele no navegador) e você será, com certeza, remetido à página verdadeira e poderá confirmar se os oferecimentos ou solicitações daquele e-mail eram verídicos.

#### 10.3.2.2. Pharming

Uma técnica de golpe bem mais elaborada que o phishing, mas com objetivo semelhante a esse. No pharming, o objetivo final é a obtenção de dados de usuários, assim como no phishing.

O modus operandi de alguém que prepara um golpe de pharming é diferente: nesse golpe, o atacante altera (adultera) as configurações de um servidor DNS, fazendo com que um domínio qualquer (como [www.bb.com.br](http://www.bb.com.br)) aponte para um endereço IP de um servidor ilegítimo, mas com um site visualmente idêntico ao do Banco do Brasil.

“Mas, João, nesse caso se eu fosse vítima do pharming não perceberia isso, porque eu mesmo pensaria ‘mas fui eu que digitei o endereço! Não há problema nisso!’, não é?”

Sim, caro leitor. Isso mesmo! Um golpe de pharming é muito mais difícil de detectar e, consequentemente, de evitar. Mas prestar atenção aos certificados de segurança dos sites

verdadeiros e conferi-los sempre quando acessar aqueles sites é algo que ajuda muito.

Essa técnica também é conhecida como DNS Poisoning (algo como “envenenamento do DNS”).

### **10.3.2.3. Engenharia social**

“Caô! 171! Lábial!”

Engenharia social é uma técnica na qual o golpista usa da persuasão, muitas vezes abusando da ingenuidade ou confiança do usuário, para obter informações que podem ser utilizadas em benefício próprio, normalmente para ter acesso não autorizado a computadores ou informações.

E quem disse que hackers são pessoas enfiadas em salas de computadores sem nenhuma capacidade de relacionamento social? Eles usam muito de psicologia e lábia para conseguir o que querem.

Só a título de exemplo: conheci um caso em que um hacker chegou a noivar com a secretária de uma empresa de telecomunicações para conseguir uma senha que lhe permitisse realizar o ataque que tanto desejava.

(Tá, eu sei, esse aí é cafajeste mesmo!)

### **10.3.3. Ataques e técnicas contra sistemas de informação**

#### **10.3.3.1 Ataques DoS (Denial of Service)**

Não é um ataque apenas, mas uma classificação de um gênero de ataques. Um ataque pertencente a esse grupo tem como característica principal o objetivo de fazer o computador-alvo parar de responder aos verdadeiros clientes que o solicitam.

Ou seja, quando ataco um computador a fim de fazê-lo travar ou desligar, de modo que, mesmo momentaneamente, ele pare de responder aos clientes que lhe solicitam, acabei de perpetrar um ataque de DoS (Denial of Service – Negação de Serviço).

Por exemplo: tirar o servidor da tomada ou então dar uma martelada na placa-mãe do servidor vítima são exemplos de ataques de DoS!

#### **10.3.3.2. Buffer Overflow (sobrecarga de Buffer)**

Uma ataque de buffer overflow (na verdade, há vários desses tipos, ou seja, é um subgênero) é realizado enviando-se mais informações do que aquelas que um determinado sistema foi programado para receber.

Se essa falha existir (ou seja, o limite de recebimento de informações existe, mas sem nenhum tipo de proteção ou garantia de evitar o fato), basta enviar uma quantidade maior de informação que a memória daquele sistema consegue entender e, dentro de pouco tempo, aqueles dados vão “passar” dos limites e invadir áreas da memória do servidor que pertencem a outras partes do programa, causando travamentos e instabilidade no servidor.

Então é só isso: sei que aquele micro não aguenta mais que 100 KB por segundo. E sei também que ninguém limitou isso (não foi previsto, pelos programadores daquele sistema, nenhum tipo de proteção para o caso de alguém tentar mandar mais que 100 KB para ele). Agora é só, simplesmente, mandar uns 300 KB e ver o “circo pegar fogo”.

Dentro de instantes o servidor irá parar, ou poderá “abrir caminho” para uma invasão, ou permitir que eu entre em um sistema de um usuário sem a senha dele, etc.

O céu (ou o “inferno”) é o limite!

#### **10.3.3.3. Ping da morte (ping of death)**

Um tipo de ataque que já foi muito comum. Esse ataque está enquadrado na classificação de buffer overflow, ou seja, ele se baseia na exploração de falhas no recebimento de dados por parte de alguns servidores no que concerne ao protocolo ICMP.

Sabemos que o PING é um comando que utiliza o protocolo ICMP. Sabemos que ele é, até certo ponto, inofensivo e, com certeza, muito útil para os administradores de rede. Mas se executarmos um comando PING para enviar pacotes com mais de 64 KB de dados para um servidor que apresente a falha, esse servidor trava, parando de responder às requisições de seus clientes legítimos.

Portanto, o mecanismo de funcionamento do ping da morte é o envio de pacotes (no protocolo ICMP, claro) com tamanhos maiores que 64 KB, pois alguns servidores ainda não sabem lidar com dados de tamanhos maiores.

Um sistema muito susceptível a ping da morte é, pasmem, o Windows 95! Esse sistema operacional simplesmente travava sempre que se mandava para ele um comando ping com tamanho maior que 64.400 bytes.

Hoje em dia, usar o ping da morte não adianta muito porque a maioria dos servidores (e sistemas operacionais de micros domésticos) atuais já foi atualizado para não ser mais vulnerável a esse tipo de ataque!

#### **10.3.3.4. SYN Flooding**

É uma técnica de ataque muito interessante que faz uso de mecanismos de conexão TCP para fazer um servidor parar de responder.

“Então também é um ataque de DoS?”

Sim, leitor! SYN Flooding também é um ataque de DoS!

Seu modo de operação se baseia no envio de pacotes SYN (sincronia), que são os pacotes iniciais para a abertura de conexão entre dois micros usando o protocolo TCP e a não realização efetiva dessa conexão.

Só para lembrar: quando dois micros iniciam um processo de comunicação via TCP, o cliente começa enviando um pacote SYN (sincronia); o servidor recebe esse pacote e envia, logo em seguida, um pacote SYN-ACK (confirmação da sincronia); depois de receber essa resposta, o cliente envia, finalmente, um pacote ACK (confirmação da conexão).

Pronto. Basta que o cliente envie vários pacotes SYN seguidamente para um servidor, que ficará de responder a todos eles e esperará a confirmação da conexão para todos eles. Ai, o cliente simplesmente não responde a nenhuma delas (não fecha as conexões).

Se todas as conexões disponíveis daquele servidor forem reservadas para aquele cliente que está perpetrando o ataque, os demais clientes (legítimos, diga-se de passagem) vão ficar “a ver navios” quando tentarem se conectar com o referido servidor. Isso se deve ao fato de um servidor possuir um número específico de conexões que pode abrir simultaneamente.

É como ter uma central de PABX com 10 ramais e alguém ligar 10 vezes e deixar a secretária esperando com os ramais ocupados. Se algum cliente real ligar para aquela empresa à procura de algum produto, vai dar de cara com o “sinal de ocupado”, pois todos os ramais estão sendo usados de forma ilegítima.

#### **10.3.3.5 Spoofing**

Não é um ataque em si, mas uma técnica usada em conjunto com qualquer ataque a ser realizado. O spoofing consiste em esconder o endereço real do atacante por meio de alteração no cabeçalho do pacote IP (IP spoofing) preenchendo-o com endereços IP falsos ou por meio de alteração do cabeçalho do quadro da rede (MAC spoofing) para que não se possa saber o endereço MAC do atacante.

Isso é usado para que não se possa descobrir, nas auditorias que são feitas após o ataque, de onde ele partiu. É uma forma de “adulterar” o endereço de remetente.

Quando faz spoofing, um atacante normalmente realiza o ataque às cegas, ou seja, sem observar qual está sendo o resultado dos seus atos, porque como os pacotes saem de seu micro com endereços de origem diferentes (fictícios muitas vezes), as respostas aos ataques nunca voltam ao atacante, que não vê o ataque efetivo, mas, no máximo, pode ter ideia do que está acontecendo.

MAC spoofing é amplamente usado em sniffers que dizem ser capazes de capturar quadros em uma rede que usa switch. Pois adulterar o endereço MAC de um micro (escrevendo endereços falsos nos quadros que vão sair pela rede) faz o switch atualizar sua tabela interna de endereços MAC com aquele novo dado.

Daquele momento em diante, qualquer quadro enviado para o endereço MAC será enviado para aquela porta específica do switch, chegando ao micro do atacante.

#### **10.3.3.6. Ataque Smurf**

Um ataque muito bem elaborado que visa tornar um micro incapaz de responder por causa de uma sobrecarga momentânea de dados (é um ataque DoS, portanto).

Para realizar um ataque smurf, usa-se o comando PING (protocolo ICMP) para enviar vários pacotes (muitos mesmo) a um endereço de broadcast qualquer (ou seja, a todos os micros de uma rede) tendo, antes do envio em si, sido alterado o endereço de origem do pacote (IP spoofing) para o IP da máquina que se deseja atacar.

Então observe. Eu mando um comando PING para todos os micros de uma rede usando, como endereço de origem deste PING, o IP de um micro qualquer (o meu alvo). O que vai acontecer?

“Todos os micros daquela rede vão responder os PING a esse pobre micro-alvo. Ele receberá tantas requisições que não conseguirá trabalhar direito!”

Exatamente, leitor!

“Que maldade, João!”

Ei! Eu só estou explicando como se faz! Não estou promovendo a realização desse tipo de coisa!

Aliás, atualmente, a grande maioria das redes locais de computadores não está mais vulnerável aos ataques smurf, pois já conseguem evitar pings para o endereço de broadcast!

### **10.3.3.7. Man-in-The-Middle (Homem no Meio)**

Técnica de espionagem e adulteração de mensagens muito bem elaborada e, se executada, muito difícil de ser detectada.

Um atacante (na verdade, um espião) usa a técnica do Man-in-the-middle (MITM) para receber mensagens de um usuário A e repassá-las para o usuário B. Em seguida, recebendo as respostas de B e repassando-as para o usuário A.

Dessa forma, A pensa que está falando com B e B pensa que está se comunicando com A, mas todas as mensagens, sem exceção, estão passando pelo bisbilhoteiro. As mensagens passadas pelo espião podem, ou não, ser alteradas (só depende do que ele quer exatamente fazer com esse golpe).

Agora chega de falar de gente ruim. Vamos falar de coisa boa! Vamos falar dos “mocinhos”, já que levamos muito tempo falando dos “bandidos”.

## **10.4. Agentes da segurança**

### **10.4.1. Antivírus**

Programa residente na memória (fica sempre na memória RAM) que protege o sistema contra infecções de vírus de computador (vírus “informático” é um nome atualmente usado) e outros malwares.

Um antivírus tanto evita novas infecções como limpa o sistema de infecções já estabelecidas. Um antivírus normalmente degrada o desempenho do computador por estar sempre executando na memória RAM e, na maioria dos casos, ser muito “pesado”. Antivírus não são sistemas efetivos contra tentativas de invasão, apenas contra malwares.

### **10.4.2. Firewall**

Programa que cria uma “barreira” de proteção contra invasores (na verdade, contra, especificamente, as tentativas de comunicação com o computador protegido). Um firewall pode bloquear as comunicações por diversos critérios, previamente estabelecidos.

#### **10.4.2.1. Filtro de pacotes**

São firewall mais simples (nossos programas firewall pessoais são assim) que normalmente atuam apenas na camada 3 (camada de rede), analisando e filtrando pacotes do protocolo IP de acordo com informações específicas contidas em seus cabeçalhos.

Como um pacote contém apenas alguns tipos de dados em seu cabeçalho (como endereço IP de origem, endereço IP de destino, porta do protocolo, entre outros), os filtros de pacotes conseguem filtrar os pacotes (decidir se passam ou são bloqueados) por meio desses poucos critérios.

Um firewall dessa categoria pode tomar decisões com base no endereço IP de origem (deixar passar ou bloquear pacotes de acordo com o endereço IP de onde vêm), no endereço IP de destino (bloquear ou deixar passar de acordo com o destino do pacote) ou ainda com base na porta do protocolo (do tipo “bloqueie todos os pacotes que venham no protocolo FTP – porta 21”).

Então, um filtro de pacotes consegue filtrar o tráfego com base em:

- a. Endereços IP de origem e destino.
- b. Porta (do protocolo) TCP ou UDP.

#### **10.4.2.2. Firewall de estado**

Os firewalls de estado (statefull firewall) são bem mais elaborados que os filtros de pacote porque trabalham na camada de transporte (analisando o tráfego TCP) e são capazes de detectar falhas não somente no nível dos pacotes (camada de redes), mas no nível das conexões TCP.

Um firewall de estado seria muito útil, por exemplo, contra um ataque do tipo SYN flooding, pois seria capaz de identificar o ataque porque analisaria a quantidade excessiva de pacotes SYN recebidos sem estabelecimento efetivo de conexão. (Um filtro de pacotes não seria capaz de identificar problemas em diversos pacotes SYN, porque não saberia ler o que são pacotes SYN – ele os deixaria passar desde que respeitassem as normas de acesso descritas na camada 3 – IPs ou portas.)

#### **10.4.2.3. Firewall de aplicação**

São filtros muito mais eficazes que os anteriores porque trabalham na camada de aplicação, analisando regras mais complexas que seus irmãos anteriores.

Esses firewalls conseguem analisar conteúdos das mensagens na camada mais alta da comunicação, sendo capazes de interagir com informações muito mais complexas e detectar potenciais problemas onde os firewalls de outros níveis não conseguem.

O único problema desse tipo de firewall é que, por ser muito complexo e cheio de recursos, ele normalmente se apresenta como um programa bastante pesado, exigindo, na maioria das casos, um computador com capacidades muito grandes para instalá-lo e usá-lo com eficiência aceitável.

#### **10.4.3. IDS**

Sistema Detector de Intrusos (IDS) é um conjunto de tecnologias (programas, hardware) que objetiva descobrir, em uma rede, os acessos não-autorizados a ela que podem indicar a ação de invasores.

IDS vindos em programas de segurança domésticos são programas que auxiliam os firewalls filtros de pacotes analisando as comunicações de uma forma mais “macro”, ou seja, ampliando a visão sobre as comunicações e alertando ao firewall sobre possíveis problemas que ele não tenha visto.

Um IDS seria capaz de detectar um ataque de SYN flooding ou a ação de um port scanner e, com isso, alertar ao firewall que bloqueie aqueles endereços IP atacantes. O firewall, sozinho, não teria condições de “ver a malícia daqueles pacotes” porque ele (o firewall) tem uma visão muito “bitolada” e só consegue analisar os cabeçalhos dos pacotes.

IDS também são sistemas que permitem a detecção da invasão enquanto já está acontecendo (com o invasor já dentro da rede).

#### **10.4.4. Antispam**



Programas que podem classificar as mensagens de e-mail recebidas como sendo aceitáveis ou como sendo spam (indesejadas). Esse programa permite que os usuários não sejam incomodados com essa prática desagradável. Como um spam pode trazer outras coisinhas chatas consigo (vírus, worms, trojans), o Antispam é um recurso bastante interessante para que nossas caixas postais sejam usadas de modo a armazenar apenas o necessário.

#### **10.4.5. DMZ – zona desmilitarizada**

Consiste em uma rede auxiliar semiprotegida, separada da rede interna da empresa, onde são hospedados os servidores daquela empresa que precisam ter acesso direto à Internet (como os servidores de páginas, de e-mail e proxies).

Em uma DMZ ficam os servidores que precisam ter acesso direto à Internet, e esses servidores têm acesso (embora restrito) aos micros da rede interna da empresa, mas os micros internos não têm acesso à Internet. Portanto, a DMZ é uma área “semiprotegida” que existe para que os servidores que precisam ter acesso à Internet não habitem junto com os computadores internos da empresa.

Com isso, os computadores internos estão em um “ambiente” mais protegido, e os serviços da Internet que a empresa oferece (e-mail, páginas, proxy) não são comprometidos.

#### **10.4.6. Bastion Host**

Um computador “superprotetor” instalado na porta de uma rede para protegê-la de todas as possíveis ameaças.

Um bastion host é um computador que funcionará como barreira para a rede, impedindo todo e qualquer ataque possível àquela rede. Um bastion host traz, instalado, um firewall (ou vários de vários tipos), antivírus, IDS etc.

Esse computador existe unicamente para proteger uma rede de computadores de ameaças externas.

#### **10.4.7. Criptografia**

Processo matemático para embaralhar uma mensagem digital, tornando sua leitura incompreensível por pessoas que não possuam a chave (código) para desembaralhar a mensagem. A criptografia pode ser usada, atualmente, para manter os dados sigilosos (privacidade) e para garantir a identidade do remetente de uma mensagem (autenticidade).

A criptografia é a “alma” dos processos de certificação digital e assinatura digital, que começaremos a estudar agora.

### **10.5. Criptografia**

Como já foi dito, a criptografia (Cripto=enigma, grafia=escrever – “A arte de escrever por enigmas”) é um processo matemático usado para embaralhar os dados de uma mensagem que deve ser secreta (confidencial).

#### **10.5.1. Entendendo a criptografia**

A principal finalidade da criptografia é, sem dúvida, reescrever uma mensagem original de uma forma que seja incompreensível, para que ela não seja lida por pessoas não autorizadas. Veja um exemplo:

|                             |                        |
|-----------------------------|------------------------|
| Mensagem original           | Mensagem embaralhada   |
| Olá, pode pagar ao cliente! | J#%9(aAs##1!2)%’&&sDoI |

A ideia só funciona, claro, se a pessoa autorizada a ler a mensagem (o receptor, destinatário, interlocutor etc.) puder transformar a mensagem embaralhada de volta em mensagem legível.

Então, temos de entender que os dois envolvidos oficiais na comunicação precisam acordar em algo. (“Acordar” não de “despertar”, claro, mas de “entrar em acordo” – essa explicação é somente para você, aluno, “acordar” desta leitura chata!)

Pense nisto:

João e José vão trocar números (senhas) pessoalmente, durante uma reunião na empresa, mas ninguém pode saber das senhas. Eles ainda não possuem tais números, mas antes da reunião, cada um vai saber qual é a sua senha.

“Ei, João, que história é essa?” – Calma! Seja criativo: imagine que João e José são espíões, sei lá!

Eles decidem (previamente, claro) que não vão passar-se mutuamente os números em voz alta, em vez disso, vão dividi-los por 43 e passar o resultado para o outro. Se fizerem isso, qualquer pessoa que não conheça o “esquema” deles será incapaz de “entender” as senhas.

A senha que João recebeu e tem de passar é 5289.

A senha que José recebeu e tem de passar é 3741.

Na reunião, João fala para José: “Ei, Zé, você viu o número de mortos no terremoto ontem? 123 pessoas!”

“Não, João” – retruca José. —“Foram apenas 87 mortos.”

José vai para casa, multiplica 123 (dado por João) por 43 (previamente acordado entre eles) e obtém 5289 (a senha).

João faz o mesmo com o 87 (dado por José) e o 43 (acordado entre eles) e obtém 3741 (a senha de José).

(Tudo bem, você vai imaginar: e se a divisão não fosse inteira – seriam 87,3 mortos? Isso é só um exemplo, não exagera nas exigências, ok?)

O que podemos tirar desse best-seller da espionagem/suspense? Simples: João e José criptografaram (cifraram) seus dados e os transmitiram em um meio inseguro, mas o meio inseguro não se apossou dos dados (ou não conseguiu entendê-los) porque eles estavam cifrados (embaralhados).

Em tempo: cifrar é o mesmo que criptografar ou encriptar. Decifrar, por sua vez, encontra sinônimos em decriptografar ou decriptar.

Mais ainda: além de saber que João e José usaram criptografia, pode-se extrair duas informações que eles previamente haviam acordado, sem as quais o processo de criptografia não pode ser completado: o Algoritmo Criptográfico e a Chave Criptográfica.

Definição: algoritmo é um conjunto finito de etapas para solucionar um problema ou realizar uma ação. Um algoritmo é um “programa” de computador, um roteiro a ser seguido pelo micro.

Algoritmo criptográfico é, portanto, o programa (ou se preferir, sequência de passos, etapas) matemático que transforma a mensagem original em mensagem cifrada, embaralhada, confusa e vice-versa.

No caso do exemplo anterior, o algoritmo foi o processo de divisão do número, afinal, João e José podiam ter escolhido coisas mais complicadas como: “primeiro, a gente divide por 10, depois a gente soma 5 e, por último, a gente multiplica por 3”. Vê que algoritmo complicado...

Se esse fosse usado, a senha de João, criptografada com esse algoritmo, seria 1601,7.

Para “decifrar” um número que passou por esse processo complicado de encriptação, as etapas deveriam ser realizadas em ordem inversa. A decriptação é, de forma bem simplória, o processo de “inverter” o que a encriptação fez. O algoritmo, em si, é o mesmo (pois as operações aritméticas são as mesmas), mas a sequência de realização delas é que é invertida.

Então, não esqueça: um dos pré-requisitos para que a criptografia aconteça é que o algoritmo seja o mesmo no processo de encriptação e no processo de decriptação.

A **chave criptográfica**, por sua vez, **é o número** que será usado, em conjunto com o algoritmo, para alterar a mensagem original. A chave é o “código” de cifragem e decifragem da mensagem. No caso do exemplo anterior, a chave era o 43 (para o processo de divisão, poder-se-ia escolher uma série de outros números).

Quanto ao uso da chave propriamente dito, a história seria um pouco diferente.

– João e José acordam que irão receber uma senha cada um e que irão trocar essa senha um com o outro. Mas, do mesmo modo, eles não vão falar a senha em voz alta, eles falarão um outro número, obtido a partir da SOMA de um valor X à senha. Em outras palavras, quando eles trocarem as senhas, eles vão informar um número X e a senha somada a esse número X. Assim:

– João recebe a senha 8526.

– José recebe a senha 5295.

– A reunião começa.

– João diz: “José, a China cresceu 450% no mercado e hoje conta com 8976 clientes no mundo.”

– José diz: “Não diga! E a Índia? Tem mais de 378 empresas de informática que atendem a 5673 clientes no mundo.”

- Ninguém na reunião entende nada do que eles disseram.
- Quando José chegar em casa vai pegar o número 8976 e o número 450, ambos dados por João, e vai fazer:  $8976-450=8526$  (essa é a senha de João).
- Quando João fizer o mesmo (ou seja,  $5673-378$ ), descobrirá que a senha de José é 5295.

Nesse exemplo, os dois “interlocutores” acordaram apenas o algoritmo criptográfico (a soma da senha com o outro número dado) e tiveram de “compartilhar” o outro número dado. Esse outro número é a chave criptográfica (se a chave mudar, o “texto cifrado” muda). Se José resolve somar para encriptar e João resolve dividir para decriptar, o resultado não será o mesmo que José enviou.

Da mesma maneira, se os dois usam chaves distintas (como José fornece a chave 546 e João resolve, arbitrariamente, usar uma chave 456), o processo de decriptação não obterá o verdadeiro valor enviado.

Então fica claro o seguinte: para funcionar corretamente, um processo de criptografia tem de usar o mesmo algoritmo criptográfico (seqüência de passos para encriptar/decriptar os dados) e, na maioria dos casos, usar a mesma chave criptográfica. (Isso é o que veremos agora.)

### 10.5.2. Criptografia é somente com números?

Bom, a princípio sim! No caso do computador, como todos os dados são digitais, eles são números (até mesmo as letras que você está lendo agora). Todos os dados que passam em um computador são digitais, portanto, são números.

Mas podemos “inventar” um código qualquer para encriptar texto. Veja o exemplo:

Algoritmo para encriptar: troque uma letra do alfabeto pela letra que estiver x posições à frente. (x é a chave).

Se eu te desse a mensagem ZEMI WXYHEV com a chave 4? O que significa?

ABCDEFGHIJKLMNOPQRSTUVWXYZ (isso ajuda?)

### 10.5.3. Criptografia simétrica (ou criptografia de chave secreta)

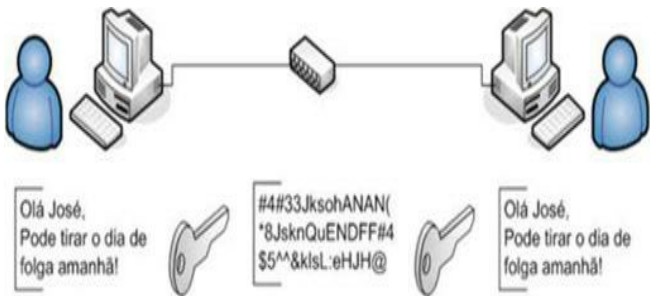
Uma das formas atuais de criptografia é chamada de criptografia simétrica, ou criptografia de chave secreta, que utiliza **uma única chave** para encriptar e decriptar os dados.

A criptografia simétrica existe há muito tempo e tem algumas características interessantes:

- a. Usa apenas uma chave para encriptar e decriptar as mensagens.
- b. É mais rápida, pois exige menos dos processadores para encriptar/decriptar as mensagens. Por esse fato, é o sistema usado para criptografar grandes quantidades de dados (como e-mails com arquivos anexos grandes ou até mesmo discos rígidos inteiros).
- c. A chave tem de ser compartilhada entre os envolvidos na comunicação, o que torna esse sistema suscetível a falhas de segurança. (Se a chave cair em mãos erradas, mensagens poderão ser lidas e forjadas pelo novo “participante da conversa”).

Os principais algoritmos de criptografia simétrica usados comercialmente por aplicações na Internet são: DES, 3DES (Triple DES) e, mais recentemente, o AES. (Todos homologados pelo Departamento de Segurança Nacional dos Estados Unidos e por normas da IEEE – e de outros órgãos – para uso comercial.) A maioria dos programas que podem fazer uso de recursos de criptografia, como os programas de e-mail, por exemplo, é capaz de usar esses algoritmos.

Veja um exemplo da criptografia simétrica:



*Figura 10.1 – Criptografia simétrica.*

Aqui vão algumas informações interessantes:

- **DES:** usa chaves de 40 e 56 bits. Já está obsoleto, mas é usado ainda.
- **3DES (Triple DES):** usa chaves de 168 bits. É três vezes mais “pesado” e mais seguro que o DES, porque criptografa a mensagem três vezes seguidas, usando DES. Então, em suma, ele usa três chaves diferentes de 56 bits que, combinadas, resultam numa chave de 168 bits.
- **AES:** usa chaves de 256 bits. É o substituto do DES e 3DES (atualmente, a grande maioria das aplicações comerciais usa esse).

A “força” de qualquer algoritmo de criptografia simétrica está ligada diretamente à chave usada (mais precisamente, ao seu tamanho) e ao processo matemático em si, utilizado por aquele algoritmo.

Os algoritmos apresentados são de domínio público, ou seja, eles já são conhecidos: muitos programas já os utilizam. Então, se os algoritmos já não são nenhum segredo, a responsabilidade pela segurança de um processo criptográfico simétrico recai sobre a chave!

#### 10.5.4. Entendendo a chave

O que é uma chave criptográfica? Como vimos, uma chave é um dado (no caso dos computadores, é um número binário) que serve de código para a encriptação e deciptação de informações.

A função da chave é simples: tornar possível, para o detentor dela, a deciptação de uma mensagem previamente encriptada. A máxima da criptografia é: “O processo de deciptação tem de ser muito simples para quem possui a chave e praticamente impossível para quem não a possui.”

Uma chave pode ser representada como algo assim: 4%5#jJhErTTIos^^)99;00sdFAAqW  
Mas, na realidade, é um número binário (porque o nosso micro só manipula sinais binários):  
01001010111010101101010100010101001110101111111010100100100100101011101010110101

#### 10.5.5. Força bruta

Então, para ser capaz de ler uma mensagem criptografada que não lhe pertence, um “invasor” tem de descobrir a chave, o que, convenhamos, é difícil, mas não impossível.

Uma chave é um número binário (uma sequência de zeros e uns). Bastaria testar todos os números binários possíveis até encontrar aquele que daria uma mensagem legível como resultado. Esses sucessivos processos de “tentativa e erro” são chamados de método de “força bruta” (brute force).

Força bruta é, em suma, programar um computador para testar, em uma mensagem criptografada, todas as possíveis chaves dentro de um espaço dado. Então, quanto mais possibilidades de chaves houver, mais difícil se torna o sucesso do método de força bruta.

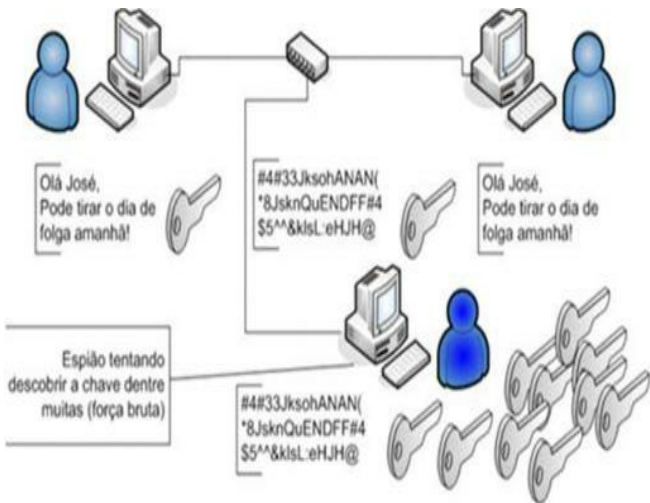


Figura 10.2 – Método de força bruta.

“Ei, João, mas também é necessário conhecer o algoritmo, não é? Quer dizer, além de descobrir a chave, é necessário saber qual o algoritmo usado para fazer a criptografia, não é?”.

Sim! Sem dúvida! Mas uma mensagem criptografada tem “assinaturas” de que algoritmo foi usado nela. Se foi um dos comercialmente usados (apresentados há pouco), os principais programas de força bruta conseguem detectá-los e usá-los, afinal, são de domínio público.

Voltemos à chave: como o algoritmo é público e todos os conhecem, fica restando ao invasor ter de descobrir a chave. Quanto mais chaves ele tiver de analisar, mais difícil e demorado será descobri-la, não é? Isso é medido pelo tamanho da chave.

#### 10.5.6. Entendendo o tamanho da chave

O tamanho de uma chave criptográfica é medido em bits (lembre-se: a chave é um número binário). Vamos começar com um processo criptográfico que utilize uma chave de 8 bits... Digamos que a minha chave é essa:

01010111

O segredo da segurança não é a minha chave em si, mas quantas possibilidades de chaves há, além da minha. Ora, oito bits são oito dígitos que podem assumir apenas dois valores (zero ou um). Por análise combinatória, teremos  $2^8$  possibilidades, ou seja, uma chave de 8 bits pode assumir 256 combinações diferentes.

Para um computador “contratado” para descobrir uma chave que é de 8 bits, será necessário testar apenas 256 chaves diferentes, e isso é muito simples, leva questão de milésimos de segundos.

Quando o tamanho da chave vai aumentando, a dificuldade vai sendo aparente: avalie uma chave com 16 bits (ainda assim, “fichinha” para os micros atuais).

São 16 dígitos binários, totalizando  $2^{16}$  possibilidades de chaves. Se um algoritmo utiliza uma chave com 16 bits, ele pode gerar 65.536 chaves diferentes. Garanto que, para um micro atual, a chave seria encontrada em menos de 1 minuto.

Hoje, são usados algoritmos, com chaves de 56 bits (DES, por exemplo) e maiores. Uma chave com tamanho de 56 bits seria algo assim:

0101001010010010101001010101010010100101001010011

Podem parecer bem frágeis, mas quantas possibilidades de chaves podem ser criadas? Claro que  $2^{56}$ , que equivalem a 72.057.594.037.927.936 (72 quatrilhões, para arredondar) chaves. As chaves de 56 bits já foram quebradas (hoje, leva-se menos de uma hora para descobrir chaves de 56 bits), portanto, elas não são mais tão seguras!

Hoje, estima-se que chaves com menos de até 100 bits são fáceis de quebrar e que as de 128 bits serão suficientemente seguras por bastante tempo ainda (não tem stress, não é mesmo? O 3DES usa chaves de 168 bits, e o AES, futuro padrão mundial, tem chaves de 256 bits).

Exemplos de utilização de criptografia simétrica podem ser vistos no dia a dia: mensagens de e-mail são criptografadas, na maioria dos casos, com algoritmos simétricos (AES é o padrão). Acessos a sites seguros (quando aparece aquele cadeado na barra de status do navegador) também são exemplos de uso de criptografia simétrica (AES normalmente).

#### 10.5.6.1. O usuário escolhe a sua chave?

Bom, isso é tecnicamente possível, mas normalmente não é assim! Quem quiser “mergulhar” no mundo da criptografia vai normalmente recorrer a um programa gerador de chaves (como o PGP, para Windows, ou o GPG para Linux).

Esse tipo de programa escolhe, aleatoriamente, uma chave para o usuário e a salva em um arquivo. A partir desse momento, essa chave poderá ser usada pelo usuário, que poderá informá-la àqueles com quem quiser se comunicar de forma sigilosa.

Aí você pergunta: “O que garante que o programa não sorteie a mesma chave para duas pessoas diferentes? E o que isso acarretaria?”

Simples! Pensa uma chave de 80 bits. Quantas possibilidades de chaves? Simplesmente 1.208.925.819.614.629.174.706.176 chaves (eu não sei nem pronunciar esse número... É 1 heptilhão, é?). Acho que é pouquíssimo provável que duas pessoas tenham a sorte de ter a mesma chave. Se isso acontecesse, elas conseguiriam ler as mensagens uma da outra, mesmo que não tivessem compartilhado suas chaves (ou seja, mesmo sem um saber a chave do outro). Agora imagine uma chave AES de 256 bits.  $1,15 \times 10^{77}$  chaves, ou, em poucas palavras, um número grande pra caramba!

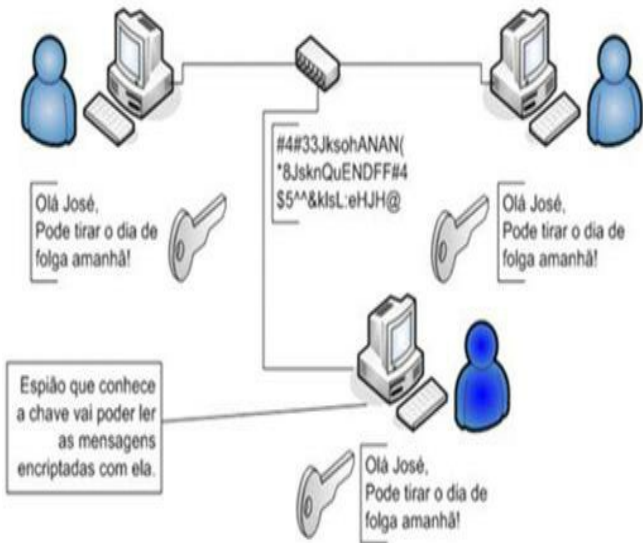
#### 10.5.7. Então, criptografia simétrica é 100% segura?

Não! Em primeiro lugar, nada é 100% seguro. O grande problema da criptografia simétrica é justamente só usar uma única chave para ambos os processos, pois: como essa chave será compartilhada?

Pense comigo um pouco: utilizamos criptografia em uma comunicação (e-mail, por exemplo) para garantir que, se houver alguém indevido bisbilhotando, ele não seja capaz de entender a mensagem, mas o destinatário devido consiga entendê-la. Certo, mas como trocar chaves com o destinatário se o meio de comunicação (e-mail) já estiver “grampeado”?

Sei lá, vamos divagar... Podemos trocar chaves fisicamente (por disquetes, CDs), mas isso requer que eu encontre cara a cara o meu interlocutor, ou envie pelo correio, o que nem sempre é possível, não acha? A chave é o segredo da segurança e, ao mesmo tempo, sua maior vulnerabilidade!





*Figura 10.3 – E a chave? Como será compartilhada entre as partes?*

Para solucionar esse problema, foi criado um novo método de criptografia: a criptografia assimétrica ou criptografia de chave pública.

#### **10.5.8. Criptografia assimétrica (criptografia de chave pública)**

Temos de divulgar a chave criptográfica para todos os envolvidos (oficiais) na comunicação, não é?

Ao mesmo tempo, temos de garantir que a chave não caia em mãos erradas, pois isso permitiria que outra pessoa (não autorizada) pudesse ler e enviar mensagens criptografadas aos demais, como se tivesse esse direito.

Informar as chaves aos envolvidos requer transferir esse dado (a chave) por um meio seguro. Mas, o que se considera seguro? E mais ainda: é seguro MESMO?

A criptografia assimétrica, também conhecida como criptografia de chave pública, veio para

resolver o problema de a chave ter de ser compartilhada (para os merecedores poderem ler as mensagens) e ao mesmo tempo evitar que os não merecedores (os bisbilhoteiros de plantão) consigam lê-las.

Na criptografia assimétrica não é criada uma única chave, mas *um par delas*. Uma das chaves serve somente para encriptar mensagens. A outra chave serve somente para deciptar mensagens.

As duas chaves são matematicamente relacionadas, não podendo haver uma delas sem a outra (ou seja, quando um programa gera um par de chaves – A e B, por exemplo –, ele não poderá gerar A sem B, nem B sem A). É como genética: as duas chaves têm o mesmo DNA, uma delas não pode ser criada em conjunto com uma terceira, só existirá com aquela outra chave especificamente. As duas chaves são geradas exatamente no mesmo momento.

“Mas por que duas chaves? E por que cada uma faz uma operação diferente?”

A chave que encripta mensagens (chamada chave de codificação criptográfica, ou chave de encriptação) será distribuída livremente, e, por isso, ela será chamada, daqui por diante, de *chave pública*, ou chave compartilhada.

Por sua vez, a chave que decipta mensagens (chamada chave de decodificação criptográfica, ou chave de deciptação) será armazenada secretamente com seu titular (dono). Essa é a *chave privada* ou chave secreta.

“Eita, João, agora embananou tudo... Por que as duas existem? E por que a privada é secreta?”

É fácil, caro:

- Eu gero um par de chaves para mim (usando um programa apropriado para isso).
- Depois de geradas, guardo, seguramente, o arquivo que contém minha chave privada.
- Envio para todos os meus amigos a minha chave pública (posso até publicar numa página da Internet para ficar mais fácil).
- Quando você quiser me enviar um e-mail sigiloso, vai usar seu programa de e-mail (Thunderbird, Microsoft Mail, por exemplo) e pedir que ele criptografe a mensagem usando a minha chave pública. Quando a mensagem for cifrada e enviada, ninguém no meio poderá decifrá-la!
- Quando o e-mail chegar à minha caixa postal, todo embaralhado, uso minha chave privada para decifrá-lo. Como os algoritmos de criptografia assimétrica definem isso: do par de chaves criado, uma delas cifra e a outra decifra (apenas isso!).

Mesmo que outra pessoa tenha a minha chave pública (o que será normal, visto que ela é “pública”), não será capaz de entender um e-mail interceptado que era direcionado a mim. Só quem poderá decifrar a mensagem sou eu, usando a minha chave privada.

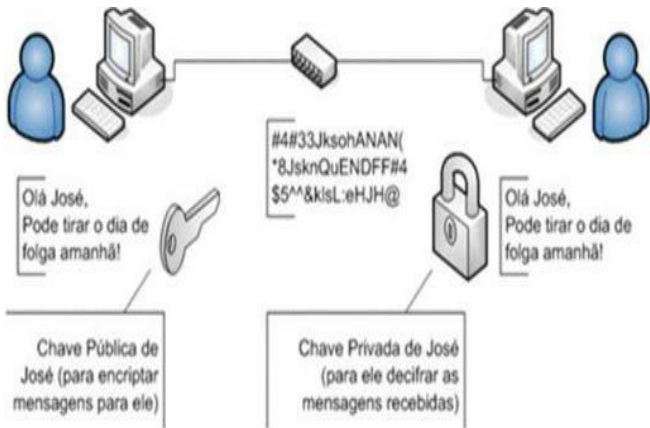


Figura 10.4 – Criptografia de chaves públicas.

“Ei, João, entendi! Mas e se você quiser mandar uma resposta sigilosa a quem te enviou o e-mail previamente? Usará que chave?”

Ao responder um e-mail de forma sigilosa, usarei a **chave pública do destinatário** para encriptar a mensagem direcionada a ele. Quando a mensagem chegar lá, ele vai usar a **chave privada dele** para decifrar a mensagem que eu envie!

Em suma: todos os usuários terão de possuir um par de chaves: uma que deverá ser mantida em segredo com cada um deles (a privada, para decifrar mensagens) e uma que deverá ser publicada, ou, pelo menos, enviada a todos aqueles de quem o usuário deseja receber mensagens sigilosas (a pública, para cifrar mensagens).

Os processos de escolha e cálculo das chaves pública e privada, ou seja, os algoritmos usados em criptografia assimétrica, utilizam métodos matematicamente difíceis de serem descobertos ou burlados, como fatoração de números primos muito grandes (da ordem de 300 dígitos, por exemplo). Pense nisso... Pense em um número primo com 300 dígitos. Agora volte rapidamente à realidade, por favor: você tem de estudar!

Dentre os algoritmos mais usados para criptografia assimétrica, está o **RSA**, usado em várias aplicações na Internet, como e-mails, páginas seguras, acesso a arquivos seguros. O algoritmo RSA utiliza chaves de 256, 512, 1.024, 2.048 e 4.096 bits.

Outra coisa que é importante saber: é possível recriar a chave pública a partir da chave

privada, pois uma é função da outra. Mas não é possível, de posse apenas da chave pública, recriar ou descobrir a chave privada associada a ela (ou seja, é uma função unidirecional – só dá para ser feita em um sentido).

Seria mais ou menos assim: qual o resto da divisão de 7 por 3? É 1, não é? Então é fácil descobrir o resultado dessa operação. Mas, no sentido inverso...

Qual é par de números que, dividindo-se um pelo outro, obtemos o resto 1? Infinitos pares dão esse resultado. Portanto, é impossível saber (descobrir), partindo do resto da divisão, o divisor e o dividendo, mas, tendo os dois, descobrir o resto é coisa simples.

Então, entenda 7 e 3 (dividendo e divisor) como a chave privada e o 1 (o resto) como a chave pública! Entendeu?

É praticamente impossível quebrar a chave privada (descobri-la) usando apenas a chave pública. Com isso, os “bisbilhoteiros” de plantão só teriam condições de “escrever” mensagens criptografadas para mim e não conseguiriam ler as mensagens escritas por outros e dirigidas a mim.

#### 10.5.9. Finalmente, uma solução segura e funcional?

Então, conseguimos uma solução para comunicação segura sem pontos negativos não é mesmo?

Não é bem assim!

Os algoritmos assimétricos são mais seguros, mas são dezenas de vezes mais lentos que os algoritmos simétricos (só para se ter uma ideia, o RSA é cerca de 1.000 vezes mais lento que o DES). E isso não seria muito indicado para anexos grandes e outras comunicações com muitos bytes (seria inviável mandar um e-mail com anexo grande criptografado em RSA porque ele demoraria muito para ser cifrado e decifrado.)

Portanto, o fato de ser muito segura e difícil de burlar torna a criptografia assimétrica algo muito “burocrático” e “pesado” para os sistemas de computação, inviabilizando seu uso em muitos casos em que a criptografia é necessária e, com isso, fazendo os algoritmos de criptografia simétrica serem usados mais comumente, mesmo com suas falhas tão “aparentes” de segurança.

Antes de prosseguirmos com a historinha, vamos resumir as principais características dos dois tipos de criptografia:

| Criptografia Simétrica   | Criptografia Assimétrica   |
|--------------------------|----------------------------|
| Usa uma única chave para | Usa chaves diferentes para |

encriptar e  
decriptar  
mensagens;

encriptar  
decriptar  
mensagens.

A chave tem  
que ser  
compartilhada  
entre os  
usuários que  
irão se  
comunicar;

Apenas  
chave c  
encriptação  
compartilhada  
(pública).  
chave c  
decriptação  
mantida e  
segredo  
(privada) co  
seu titular.

Existe apenas  
uma única  
chave para  
todos os

Cada usuário  
que irá s  
comunicar  
possui um p

envolvidos na  
comunicação.

Os processos  
de encriptação  
e deciptação  
são simples  
(exigem pouco  
processamento)  
– ideal para  
grandes  
quantidades de  
dados.

É mais  
suscetível a  
quebras de  
segredo da  
chave; Ataques  
de força bruta

de chave  
próprio.

Os processos  
são mais lentos  
(exigem mais  
cálculos de  
processadores;  
– viáveis  
apenas em  
pequenas  
quantidades de  
dados.

É praticamente  
impossível  
quebrar as  
chaves atuais  
em tempo  
suficientemente

são a mais indicada forma de quebrar a chave (descobri-la).

Principais algoritmos:  
DES: Chaves de 40 e 56 bits  
3DES: Chaves de 168 bits  
AES: Chaves de 256 bits

hábil (ne mesmo usando vários computadores reunidos).

Principal algoritmo:  
RSA: Chaves de 256 bits, 512, 1024 e até 4096 bits

Mas, para consolo dos mais neuróticos, há soluções, usadas hoje em dia na maioria dos casos, que aliam as melhores características dos dois tipos de criptografia em uma única criptografia “híbrida”.

#### 10.5.10. Criptografias simétrica e assimétrica juntas: um exemplo simples

Pense que o usuário João pretende enviar um e-mail, com um arquivo do Word anexo a ele, de forma sigilosa para José. Que criptografia usar? Simétrica porque é mais rápida (demoraria menos tempo e João não pode esperar) ou assimétrica porque é mais segura (usando, para isso, a chave pública de José), garantindo que ninguém mais, além do próprio José, será capaz de decifrar a mensagem?

Que tal usar as duas? Entenda comigo: João escreve a mensagem e anexa a ela o arquivo do

Word e, faz o seguinte:

1. João criptografa a mensagem usando um algoritmo simétrico, digamos, AES com uma chave de 256 bits, porque será mais rápido criptografar a mensagem e o arquivo do que com uma chave assimétrica. Essa chave será criada aleatoriamente somente naquele momento, para aquela transação somente (chamada “chave de sessão”) e será descartada tão logo José leia a mensagem do outro lado.
2. Mas o problema é como enviar essa chave de sessão a José com a garantia de que não haverá ninguém bisbilhotando. Simples: criptografe essa chave usando a chave pública de José. Isso significa que a chave de sessão só será “descoberta” pela chave privada de José, ou seja, apenas por ele.
3. A mensagem de João, enquanto trafegando pela rede, será composta do conteúdo original da mensagem, cifrado com a chave de sessão (a chave simétrica aleatória) e da própria chave de sessão cifrada com a chave pública de José.
4. Quando a mensagem for recebida por José, ele usará sua chave privada para decifrar a chave de sessão.
5. Depois disso, José usará a chave de sessão para decifrar a mensagem em si, podendo, finalmente, lê-la de forma clara (inclusive seu anexo).



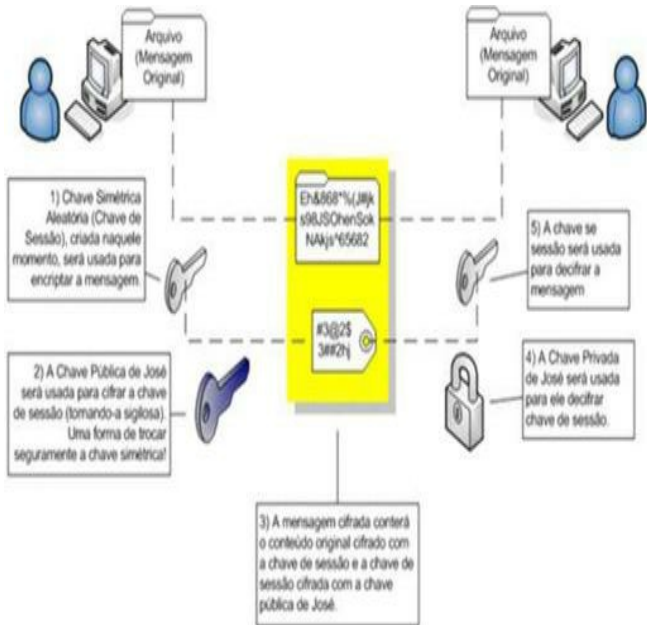


Figura 10.5 – Criptografias simétrica e assimétrica juntas.

“É muito complicado fazer isso tudo, não?”

Não se preocupe, leitor! É tudo transparente, invisível ao usuário: praticamente todos os programas de envio e recebimento de e-mail são capazes de realizar essa “operação complicada” sem que o usuário perceba. A mensagem será transferida com essas duas encriptações e com a chave de sessão em seu interior.

#### 10.5.11. A criptografia garante o quê?

Como a criptografia tem como intuito fazer com que uma mensagem não seja lida por pessoas

não autorizadas, o princípio da segurança atingido por essa técnica é, sem dúvidas, a **confidencialidade** (sigilo).

A criptografia não garante a integridade dos dados, porque eles podem ser alterados durante uma interceptação. Essa alteração pode até não ter sentido, visto que não necessariamente o espião saberá o que está fazendo (quando ele não conhece a chave), mas em alguns casos (quando conhece/descobre a chave), a alteração poderá ser realizada no meio do caminho e o destinatário não conseguirá detectar alterações.

Outra coisa que a criptografia não faz é garantir a identidade do remetente de uma mensagem. Por exemplo, se João e Jorge possuírem a chave pública de José, Jorge poderá enviar um e-mail criptografado a José, forjando os dados de envio e fazendo parecer que foi João que o fez. Portanto, a criptografia (sozinha) não garante a autenticidade.

Se não se pode garantir a identidade de um usuário remetente, não é possível garantir o não repúdio. Portanto, o usuário poderá, normalmente, negar que foi ele que enviou tal mensagem (mesmo tendo sido ele).

## **10.6. Resumo da Mensagem (Message Digest) – Hash**

Há um método matemático bastante usado para garantir a integridade dos dados durante uma transferência qualquer (ou seja: garantir que o dado não foi alterado no meio do caminho). Esse recurso é conhecido como hash (fala-se “Résh”), ou Resumo da Mensagem (Message Digest, em inglês).

O hash é uma função matemática unidirecional (pode ser feita em um sentido e não no outro – como a relação entre as chaves em um sistema de criptografia assimétrica) para escrever uma quantidade definida de bytes relacionada a uma mensagem de qualquer tamanho. Em suma: não se pode obter o arquivo original a partir do hash (é impossível, ou, no mínimo, bastante improvável).

O hash é como o dígito verificador do CPF, que está lá para confirmar a sequência de números que o antecede (o CPF propriamente dito).

Não importando o tamanho da mensagem ou arquivo, ele terá sempre um hash de tamanho fixo (isso depende unicamente do algoritmo de hash utilizado).

### Pérolas da Poesia

"Minha terra tem palmeiras  
Onde canta o Sabiá  
As aves que aqui gorjeiam,  
Não gorjeiam como lá!"

"Subi num pé de manga pra pegar abacaxi,  
Como não é tempo de morango,  
Roubaram minha bicicleta"

Rs@A221HhjErr\$%mNDDp

### Resumo da reunião:

- Implementar a segurança física e lógica;
- Comprar equipamentos necessários;
- Realizar treinamento nos funcionários;
- Definir metas de segurança de acesso;
- Estabelecer limites aos usuários comuns;
- Impedir o acesso à rede da empresa de casa;

Data: 26/11/2005

TrHa@@2!15\*&87FFg\$Kl

*Figura 10.6 – Exemplo de mensagens com hashes.*

Como o hash funciona na prática? Vamos exemplificar como ele vai funcionar para e-mail, por exemplo:

1. Quando um e-mail é enviado, o remetente calcula o hash da mensagem, que dá um resultado com tamanho definido (digamos, 20 caracteres) como: Asd#234iOO9Qne\$KELd@.
2. O remetente, então, envia o hash junto com a mensagem.
3. Quando a mensagem chega ao destinatário, este também calcula o hash da mensagem e o compara com o hash enviado pelo remetente.
4. Se o resultado do cálculo do destinatário apresentar um valor idêntico ao do hash enviado do remetente, garante-se a integridade dos dados enviados (ou seja, eles não foram alterados durante o percurso remetente-destinatário).

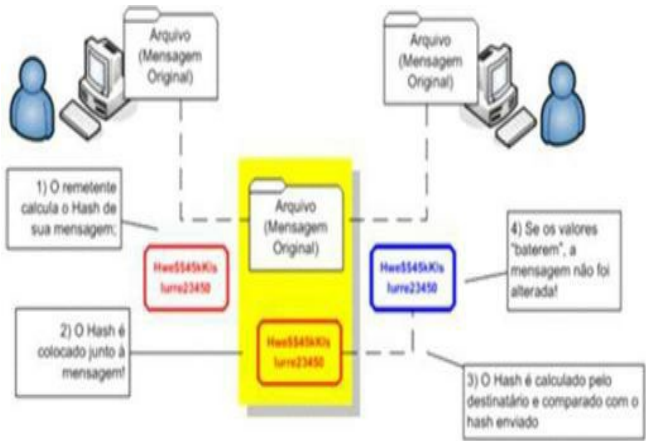


Figura 10.7 – Envio de e-mail com hash.

Essa “certeza” de que não houve alterações se dá por um simples motivo: é praticamente impossível que duas mensagens diferentes apresentem o mesmo hash!

Além disso, o hash é incrivelmente sensível a alterações. Em um nível absurdo em que se for colocado um espaço em branco a mais em um texto (um simples espaço), o hash tende a ser completamente diferente.

O hash é muito utilizado por pessoas que baixam da Internet arquivos muito grandes. (O pessoal do Linux conhece bem porque os arquivos enormes do Linux já são postos nos servidores acompanhados de arquivos de texto com o resultado do cálculo de hash.)

Veja um exemplo: a figura a seguir mostra alguns arquivos num servidor de FTP do BrOffice.org. Ao lado, a lista de hashes desses arquivos (que servirá para o usuário que os baixar ter certeza de que o download ocorreu perfeitamente). O algoritmo usado no resumo daquela mensagem é o MD5 (muito usado pela comunidade de usuários do Linux).

|                                                           |                                                           |
|-----------------------------------------------------------|-----------------------------------------------------------|
| Parent Directory                                          | BrOffice 3.2.1 Linux x86-64 install-deb pt-BR.tar.gz      |
| BrOffice 3.2.1 Linux x86-64 install-rpm-v39E pt-BR.tar.gz | BrOffice 3.2.1 Linux x86-64 install-rpm-v39E pt-BR.tar.gz |
| BrOffice 3.2.1 Linux x86 install-deb pt-BR.tar.gz         | BrOffice 3.2.1 Linux x86 install-rpm pt-BR.tar.gz         |
| BrOffice 3.2.1 Linux x86 install-rpm pt-BR.tar.gz         | BrOffice 3.2.1 Linux x86 install-rpm-v39E pt-BR.tar.gz    |
| BrOffice 3.2.1 MacOS x86 install pt-BR.dmg                | BrOffice 3.2.1 Win x86 install pt-BR.exe                  |
| BrOffice 3.2.1 Win x86 install-v39E pt-BR.exe             |                                                           |

Figura 10.8 – Hashes em MD5 dos arquivos do BrOffice.org.

Os algoritmos de hash mais comuns são:

- **MD4 e MD5:** criam um resumo de 148 bits (16 caracteres).
- **SHA-1:** cria um resumo de 160 bits (20 caracteres) – é o mais usado atualmente.

Você pode estar pensando: “João, como é que o algoritmo MD5 usa 16 caracteres, segundo o que você diz, mas no exemplo da figura, esse hash tem 32 caracteres?”

Fiquei me perguntando isso também por alguns minutos, mas aqui vai a resposta: o hash apresentado na figura está representado em hexadecimal (ou seja, só usa os caracteres 0,1,2,3,4,5,6,7,8,9,A,B,C,D,E e F). Como cada dígito hexadecimal é composto por 4 bits, o hash da figura tem, na verdade, 148 bits mesmo, que, quando convertidos em ASCII – o código de caracteres mais usado na informática – vai passar a ser representado por 16 caracteres (entre letras, números, símbolos, sinais de pontuação e tudo mais que se pode digitar). Quando contei os caracteres do MD5 (16), me referi aos caracteres ASCII (ou seja, letras, números e símbolos que podemos digitar).

### 10.6.1. As famílias de algoritmos de hash

Durante muito tempo, a família de algoritmos MD (Message Digest) foi usada comercialmente para fazer hash. Já existiu o MD1, MD2, MD3 e os últimos integrantes famosos foram o MD4 e MD5. Devido a uma série de problemas nesses algoritmos, eles foram substituídos pela família SHA (Secure Hash Algorithm).

O mais famoso membro da família SHA é o SHA-1, usado em muitos casos, inclusive em assinaturas digitais, acessos a bancos na Web, entre outros. Já existem outros mais novos como o SHA-224, SHA-256, SHA-384 e SHA-512 (esse grupo recebe o nome, normalmente, de SHA-2).

Já foram descobertas falhas nos MD (por isso foram substituídos) e, recentemente, falhas no SHA-1. Não foram, ainda, descobertas falhas na família SHA-2.

“Que tipo de falhas, João?”

Colisão de hashes, por exemplo.

“Como assim, João? O que é uma colisão de hashes?”

Simple: o hash tem algumas regrinhas básicas:

- a. Deve-se ser impossível de encontrar a mensagem original partindo-se da análise do hash apenas (isso todo algoritmo consegue!).
- b. O hash tem de parecer aleatório, mesmo que todo mundo conheça o algoritmo. Ou seja, qualquer mudança mínima na mensagem (uma vírgula colocada em um local qualquer do texto) tem de gerar um hash completamente diferente do hash da mensagem anterior.
- c. Deve ser impossível encontrar duas mensagens com o mesmo hash.

Se duas mensagens diferentes apresentam o mesmo hash, há uma colisão de hashes, o que, convenhamos, não é impossível, porque as mensagens são infinitas, mas os hashes não! Normalmente quando uma colisão é descoberta, um algoritmo é descartado e passa-se a utilizar outro mais forte (mais aleatório e com menos chances de gerar colisões).

Tem mais: o fato de o hash parecer aleatório garante que não há possibilidade de alguém forjar um hash. Quer dizer, conseguir uma colisão não é algo proposital, mas aleatório! Não se forja uma mensagem com hash igual ao de outra. Se isso acontecer, foi pura sorte!

Por esse fato, tem-se o hash como algo seguro: um espião não pode alterar um e-mail de modo que a mensagem alterada dê como resultado o mesmo hash da mensagem original; se ele conseguir isso, é melhor parar de espionar (pois dá trabalho e não dá tanto dinheiro) e jogar na mega-sena!

#### 10.6.2. O que obtemos com o hash?

Bom, com o hash atinge-se, sem dúvidas, a garantia de **integridade** dos dados transmitidos. Mas somente isso!

Como o hash não criptografa a mensagem, não se consegue confidencialidade. Como não há garantias de quem mandou a mensagem (porque qualquer um pode ter calculado o hash antes de enviá-la), não há autenticidade (e, com isso, não há garantia de não repúdio).

Quer entender mais sobre hash e criptografia? Visite [www.euvoupassar.com.br](http://www.euvoupassar.com.br) e assista ao curso de Segurança da Informação, ministrado por mim!

#### 10.7. Assinatura Digital

Já fomos apresentados a um recurso que garante o sigilo (confidencialidade) dos dados (a criptografia) e a outro recurso que garante a integridade dos dados (o resumo da mensagem, ou hash).

Agora é hora de conhecermos um recurso da comunicação digital que garante a autenticidade dos dados, permitindo associar um determinado dado a um determinado usuário remetente: esse recurso é a **Assinatura Digital**.

A assinatura digital se baseia em criptografia assimétrica, ou seja, na existência de um par de chaves para cada usuário (uma pública e outra privada). A principal diferença entre a criptografia assimétrica e a assinatura digital é como essas chaves serão usadas.

No processo criptográfico, no intuito de se ter confidencialidade (sigilo), o remetente usa a chave pública do destinatário para encriptar a mensagem esperando que ele (o destinatário) seja capaz de decifrar a mensagem usando a chave privada dele (destinatário). Então, em suma, a comunicação sigilosa usa apenas as chaves do destinatário da mensagem.

No processo de assinatura digital, com o qual se deseja a autenticidade, o remetente usará sua chave privada para “assinar” a mensagem. Do outro lado, o destinatário usará a chave pública do remetente para confirmar que ela foi enviada realmente por aquela pessoa. A mensagem não é sigilosa, porque não é criptografada, e também porque teoricamente “todos” têm acesso à chave pública do remetente (afinal, ela é pública).

A assinatura em si é apenas um conjunto de dados colocados junto à mensagem mediante um cálculo matemático feito com a chave privada do remetente em relação àquela mensagem. Vamos a um exemplo mais prático: em um e-mail, uma assinatura é apresentada como um “arquivo anexo” à mensagem e esse arquivo traz um resumo daquela mensagem devidamente processado (encriptado) pela chave privada do remetente.

O destinatário confirmará que foi o remetente que a enviou fazendo um cálculo (decriptando aquela assinatura) que atestará aquele resultado da assinatura (que, por sinal, matematicamente só poderia ter saído daquela chave privada).

Em suma, a assinatura digital utiliza apenas as chaves do remetente para a comunicação, diferente da criptografia em si.

Em poucos passos, a assinatura digital funciona da seguinte maneira:

1. O remetente (João) escreve um e-mail para José e o assina, usando, para isso, sua chave privada.
2. A mensagem não é criptografada, é apenas “assinada”, ou seja, o remetente usa sua chave privada para gerar um valor numérico associado ao e-mail. Esse valor é único para cada chave (ou seja, não há duas pessoas no planeta – ou será muito difícil – que tenham mesma assinatura).
3. Quando a mensagem chegar ao destino, José usará a chave pública de João (óbvio que José tem de possuí-la antes de qualquer outra coisa) para confirmar, matematicamente, que se trata da assinatura feita pela chave de João.
4. Quando confirmada, José pode ter certeza de que a mensagem de e-mail realmente veio de João (pois é muito difícil que outra pessoa possua sua chave privada, a menos que ela – a chave – tenha sido extraviada).

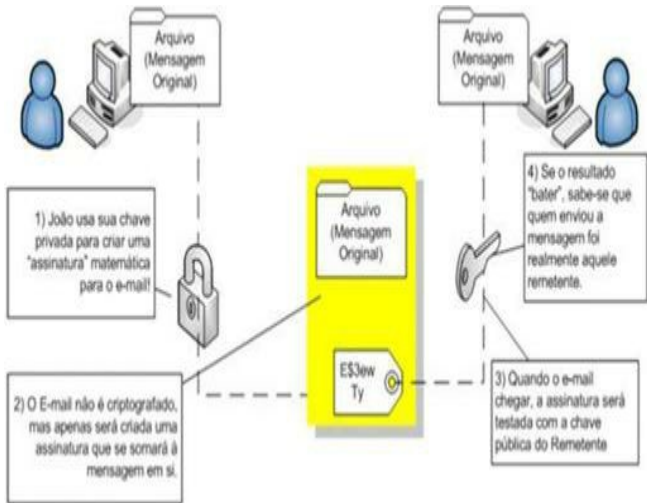


Figura 10.9 – Assinatura digital.

A grande maioria dos programas de correio eletrônico da atualidade é capaz de manipular mensagens assinadas e, é claro, assiná-las também.

Da mesma forma que existem algoritmos para hash e criptografia, há também algoritmos para assinatura digital, que também é um processo matemático. Dentre os diversos, podemos citar o **DSA** (Digital Signature Algorithm).

O algoritmo DSA usa as chaves criptográficas de um usuário para assinar mensagens. Esse algoritmo não pode ser usado para criptografia, somente para assinatura digital.

#### 10.7.1. Assinatura digital na prática

A teoria a respeito da assinatura digital é muito interessante porque se garante, com ela, a autenticidade de um usuário, devido ao fato de a mensagem ter sido assinada pela chave privada de tal usuário.

Mas, na prática, a assinatura digital é mais que isso: ela não só garante a autenticidade, mas a integridade dos dados enviados (ou seja, a assinatura digital não só garante que foi João que



mandou a mensagem, ela também garante que João mandou exatamente aquela mensagem, ou seja, ela não foi alterada no meio do caminho).

Isso significa que, para uma mensagem assinada, João não pode dizer “Ei! Não fui eu que escrevi essa mensagem!” e também não pode dizer: “Eu escrevi, mas ela não era assim! Foi adulterada!”. Isso significa que a assinatura digital garante a condição de não-repúdio!

“Mas, por que, João?”

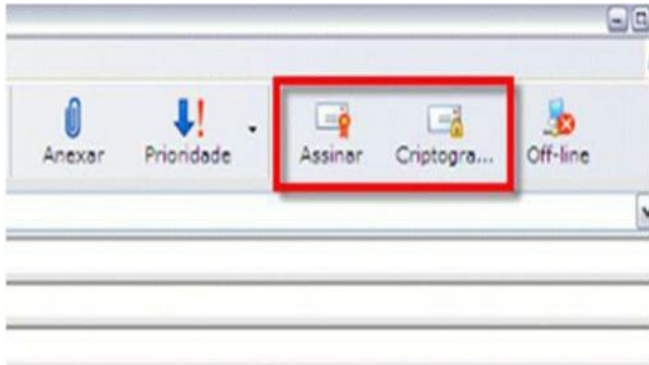
Porque a assinatura digital usa o recurso de hash para assinar a mensagem, tornando-a íntegra, autêntica e fazendo isso sem utilizar de muito processamento (faz rapidamente).

Como funciona mesmo?

1. João escreve um e-mail.
2. João calcula o hash da mensagem (usando SHA-1, por exemplo).
3. João criptografa o hash da mensagem usando DSA (isso é a assinatura) com sua chave privada.
4. A mensagem é enviada.
5. José recebe a mensagem.
6. José calcula o hash da mensagem.
7. José decifra o hash enviado com a chave pública de João e o mesmo algoritmo (DSA, no caso).
8. José compara o hash calculado por ele com o hash enviado por João. Se os dois forem iguais, então a mensagem está íntegra e realmente foi enviada por João. Se houver alguma diferença (mesmo que mínima) entre os dois hashes, então não se pode garantir a integridade nem a autenticidade da mensagem.

Lembre-se: o cálculo do hash e a assinatura são processos realizados pelos programas de correio de João e José. Ou seja, os usuários não têm de ficar calculando nada com lápis e papel.

Quando você recebe um e-mail assinado, clica em um botão para comparar as assinaturas. Se baterem, a mensagem está ok: íntegra e autêntica! Veja, a seguir, uma foto do Outlook Express e os botões de criptografia e assinatura digital do programa:



*Figura 10.10 – Botões de criptografia e assinatura digital.*

### 10.7.2. O que se obtém com a assinatura digital?

Bem, com a assinatura digital pode-se garantir:

- **Autenticidade:** o fato de a assinatura ter sido realizada pela chave privada do remetente e confirmada por sua chave pública (no destino) oferece a garantia de que foi realmente aquele usuário que a enviou.
- **Integridade:** como a assinatura digital usa hash, é possível garantir que a mensagem não foi alterada no meio do caminho.

E, com essas duas...

- **Não Repúdio:** o usuário não poderá dizer que não foi ele quem escreveu aquela mensagem.

Com a assinatura digital não se consegue a confidencialidade porque não há criptografia dos dados, que seguem “abertos” pelo e-mail para qualquer um ver, além disso, qualquer um que possua a chave pública de João conseguirá, ao receber o e-mail, verificar a assinatura e confirmar que o e-mail veio realmente de João.

Com a assinatura digital, é possível associar, de forma unívoca, um documento digital a uma chave privada e, consequentemente, a um usuário. A assinatura digital é o processo que baseia a validade jurídica de documentos digitais.

### 10.7.3. A assinatura digital serve sozinha?

“Então, João, a assinatura digital é a solução para todas as questões relacionadas à validade

jurídica de um documento?”

Não, nobre leitor. Sozinha, a assinatura digital é apenas um ato de atrelar um documento, ou uma operação a uma chave privada. Mas essa “relação” entre a chave e o documento só será válida perante a justiça se houver uma “institucionalização” das chaves envolvidas, ou seja, se as chaves usadas no processo de assinatura forem “oficiais”.

“Como assim?”

Veja bem, caro leitor, você alugaria um apartamento seu, ou venderia um carro, ou prestaria um serviço pago para alguém que não conhece? E se ele assinar um contrato, você faria? E se ele assinar um contrato com firma reconhecida em cartório, você se sentiria mais seguro?

A palavra de ordem é: **confiança** (novamente!).

Se você confia no indivíduo, ele será sua “garantia” de que cumprirá o contrato. Se você confia que ele honrará a assinatura, ela (a assinatura no contrato) é a sua garantia de que você poderá exigir depois (embora ele possa repudiá-la). Mas, com certeza, confiar assim é confiar quase que cegamente.

Quando a assinatura está “atestada” por um terceiro em quem as duas partes confiam, tem-se um nível satisfatório de confiança. Então, somente dá para se “encostar” no travesseiro e dormir tranquilo quando se confia nas partes e/ou no terceiro que certifica a idoneidade (ou, pelo menos, a identidade) das partes.

Nesse caso, entra o cartório que reconheceu a assinatura no contrato e deu seu carimbo (normalmente, agora, um selo impresso eletronicamente) – ele é o terceiro de confiança. Ele (o cartório) é o componente no qual as duas partes confiam e que certifica a identidade de uma para a outra.

Desconfiar do cartório é jogar por terra todo o alicerce dessa estrutura: a confiança.

No caso de transações pela Internet, a assinatura digital tem tido uma grande importância para confirmar a autenticidade de mensagens transferidas pela Grande Rede, mas ela sozinha não é nada, porque se pode “forjar” uma assinatura, criando um par de chaves aleatórias para esse fim. Pense comigo:

- a. Você recebe um e-mail assinado por João, com quem troca mensagens já há algum tempo.
- b. Quando você pede para confirmar a assinatura digital da mensagem, vê que realmente foi João quem mandou a mensagem. Isso porque, você compara o hash da mensagem e o decripta usando a chave pública de João.
- c. Mas, aqui vai a pergunta que não quer calar: mesmo se seu computador “achar” que a mensagem foi enviada por João e mesmo que você tenha recebido aquela chave pública achando que é de João, você pode garantir que do outro lado existe mesmo um João? Como saber se ele (João) não é outra pessoa? Como saber se o João não passa de um personagem na cabeça doentia de um maniaco homicida? (Que drama!)

Então, chegamos ao ponto em que a Justiça, e o governo como um todo, encontra falhas no uso da assinatura digital com chaves quaisquer para que os documentos digitais tenham validade e possam não ser repudiados. Mas foi essa falta de “certeza” (ou melhor, falta de confiança no interlocutor) na Internet que demandou a criação de um modelo de “confiança” coletiva conhecido como certificação digital.

## 10.8. Certificação Digital

A certificação digital é um processo que garante, de forma única, a identidade de uma pessoa (usuário de e-mail, por exemplo), ou de um computador (quando acessamos o banco). A certificação digital é garantida por um terceiro de confiança: uma instituição conhecida, normalmente, como AC (Autoridade Certificadora – CA em inglês). A certificação digital se baseia na existência de documentos chamados Certificados Digitais para cada indivíduo a ser autenticado (pessoa ou micro).

Um certificado digital é um documento (um arquivo em seu computador, por exemplo) que guarda informações sobre seu titular e é atestado (garantido) por uma autoridade certificadora.

Dentre os dados contidos no certificado digital, podemos citar:

- a. O nome completo do titular do certificado.
- b. O endereço de e-mail do titular do certificado (se necessário).
- c. A **chave pública** do titular do certificado (**obrigatório**).
- d. O nome da autoridade certificadora.
- e. A assinatura da autoridade certificadora (é isso, e a confiança na autoridade certificadora, que faz o certificado ter validade).
- e. Algumas informações adicionais (isso depende da necessidade: endereço, CPF, Identidade, Título de Eleitor, PIS/PASEP etc.).

Um certificado é, em poucas palavras, a nossa chave pública, somada a mais alguns dados importantes, assinada pela autoridade certificadora! Sim! Uma chave pública que foi assinada digitalmente pela AC! Só isso!

Existem certificados para várias finalidades, como: provar a identidade de um remetente de e-mail (autenticidade de e-mail), provar a identidade de um servidor (computador) com quem nos comunicamos (autenticidade de servidor) ou provar a nossa identidade para um site qualquer quando esse nos requisita uma identidade digital (autenticação de cliente).

Quando mandamos um e-mail assinado com nosso certificado, ele chega ao destinatário e este se encarrega de “checar” as nossas credenciais junto a quem emitiu o certificado: a autoridade certificadora. Mas é tudo automático: quando o e-mail chega assinado digitalmente, clicamos no ícone correto e temos chance de ver o certificado e, com isso, confiar nele (o certificado) e, com isso, confiar na mensagem.

Veja, a seguir, um exemplo de um certificado salvo em um computador com Windows (o arquivo HSBC.p7b é o certificado do site [www.hsbc.com.br](http://www.hsbc.com.br)):



*Figura 10.11 – Certificado digital salvo no computador.*

“Quem emite o certificado? Quem atesta que ele é autêntico?”

Em ambos os casos, a resposta é AC (Autoridade Certificadora). A AC é a instituição (privada, normalmente) que funciona como um “cartório” na Internet: é a AC que emite certificados e é a AC que confirma sua validade e autenticidade quando for consultado.

Como um certificado funciona? Vamos imaginar uma coisa bem simples: o acesso a um site da Internet que é considerado “seguro”. Você sabia que o site que você está acessando pode ter sido forjado? Sim! O site do banco que você está acessando, e no qual você digita a sua agência, conta-corrente e senha, pode, em alguns casos, não pertencer ao seu banco realmente. Como saber então? Pelo certificado.

“Ei, não preciso de certificado. Sei que um site é seguro e autêntico pelo *cadeadinho* que aparece na parte de baixo da janela do programa!”

Doce ilusão, nobre leitor... Doce ilusão...

Muita gente acha que pode garantir a autenticidade de um site simplesmente pela presença do cadeado (o ícone) na barra de status do programa navegador (ou na barra de endereços, hoje em dia), mas isso não é uma garantia da autenticidade do site, é uma garantia apenas de que a comunicação entre o site e o usuário está sendo feita de forma criptografada.

Ou seja, o cadeado indica que a comunicação entre o seu programa navegador e o site que você está acessando está sendo feita via HTTPS (o que também pode ser constatado no próprio endereço).



Então, em suma, o cadeado não garante a autenticidade do site com quem se está conversando, mas garante a confidencialidade dos dados dessa conversa (ou seja, se houver alguém “bisbilhotando” a conversa, não conseguirá entendê-la).

Essa “conversa sigilosa” pode ser forjada também, quer dizer, um site “clandestino”, que finge ser um site de um banco, por exemplo, pode criar uma conexão criptografada somente para que o cadeado apareça e o usuário acredite se tratar realmente do site do seu banco.

É aí que entra o certificado: quando você entrar no site “seguro” e visualizar o cadeado, acione um clique duplo no cadeado: ele irá abrir uma janela mostrando as “regras” dessa comunicação segura, inclusive, mostrando acesso ao certificado daquele site, como mostrado na figura a seguir.

Geral

Detalhes

Caminho de Certificação

**Informações sobre o Certificado****Este certificado destina-se ao(s) seguinte(s) fim(ns):**

- Garante a identidade de um computador remoto
- Prova a sua identidade para um computador remoto
- 2.16.840.1.113733.1.7.23.6

\* Veja a declaração da autoridade de certificação para obter d

**Emitido para:** www3.hsbc.com.br**Emitido por:** VeriSign Class 3 Extended Validation SSL SGC CA**Válido a partir de** 19/ 09/ 2012 **até** 30/ 09/ 2014

Declaração do Emissor

Saiba mais sobre [certificados](#)

OK

Note que o certificado foi emitido pela Verisign (uma Autoridade Certificadora bastante conhecida – talvez o “mais confiável” cartório virtual da Internet). Esse é um certificado que autentica a identidade do servidor (no caso, o servidor do HSBC). Esse certificado é a garantia de que, quando acessei o endereço do HSBC, a página que eu estava vendo realmente era do HSBC, e não de algum “forjador” de sites.

Se, ao dar duplo clique no cadeado, o seu navegador exibir uma mensagem do tipo “Este certificado não é válido, ou expirou, ou foi emitido por uma Autoridade Certificadora em quem você não confia”, desconfie do site, desconfie da autenticidade do site e, pelo amor de Deus, não compre ou insira dados sigilosos de maneira alguma!

Da mesma forma, quando se envia um e-mail assinado digitalmente com um certificado, o destinatário tem condições de analisar esse certificado digital, que, se for emitido por uma autoridade certificadora confiável, será suficiente para provar, ao destinatário, a identidade do remetente e a integridade da mensagem.

### **10.8.1. Validade do certificado**

Um certificado, como pode ser visto na figura anterior, que mostra seu conteúdo, tem validade (data para expirar). Depois de expirado um certificado, é necessário solicitar sua renovação para que ele possa continuar sendo usado pelo usuário, empresa ou computador.

Quando a renovação acontece, o arquivo usado anteriormente passa a não ser mais necessário, e um novo arquivo (um novo certificado) é emitido para o usuário, usando a mesma chave pública que ele detinha antes. Ou seja, o usuário em questão não ganha um novo par de chaves gerado, ele usa o mesmo par, mas a chave pública é assinada novamente pela AC e essa assinatura tem nova data de validade.

Um certificado pode perder a validade antes do prazo, se for solicitada, pelo seu titular, a sua revogação. A revogação acontece, normalmente, quando o certificado é roubado, extraviado, perdido ou quando se desconfia que ele está sendo usado por outra pessoa.

Quando um certificado é revogado, ele é colocado em uma listagem (Chamada LCR – ou Lista de Certificados Revogados – ou RCL, em inglês) e publicando no site da AC na Internet.

### **10.8.2. Analisando um certificado (problemas que podem ocorrer)**

Quando um computador recebe um certificado digital (seja no acesso ao site de um banco, seja no recebimento de um e-mail assinado, não importa), ele analisa algumas “coisinhas”. Vamos, por exemplo, imaginar que acessamos um site que nos enviou seu certificado. Nosso programa navegador vai avaliar, basicamente, quatro questões.

- **Data de Validade do Certificado:** se um certificado apresentado ao nosso navegador estiver com a data de validade vencida, o nosso programa vai avisar em uma caixinha de diálogo: Este certificado está expirado! Deseja acessar o site mesmo assim?
- **Revogação:** quando o certificado é apresentado, o nosso programa analisa se ele faz parte da lista de certificados revogados no site da AC que o emitiu. Se esse certificado estiver revogado, o nosso navegador avisará com todo prazer: Certificado Revogado! Deseja



continuar?

- **Titular do Certificado:** eis um problemão – se o certificado apresentado foi emitido para um site diferente daquele que está apresentando o certificado, aí temos algo de que desconfiar!

Imagine que você acessa o site [www.qualquercoisa.com.br](http://www.qualquercoisa.com.br) e esse site envia para você um certificado emitido para [www.outracoisa.com.br](http://www.outracoisa.com.br).

Isso é deveras estranho! O navegador vai dizer em letras garrafais: O certificado apresentado por esse remetente não foi emitido para ele. Deseja continuar mesmo assim?!

- **Confiança no Órgão Emissor (AC):** em alguns casos, entramos em sites confiáveis, com certificados legítimos, mas que foram emitidos por instituições nas quais não confiamos. Sim! Nós é que escolhemos em quem confiar! A lista de autoridades certificadoras nas quais confiamos está presente em nosso navegador (nas configurações dele).

Se o certificado do site foi emitido por alguma AC que não está presente na nossa lista de confiança, o browser vai reclamar: O certificado deste site foi emitido por uma instituição na qual você não confia. Deseja acessar o site assim mesmo?

Técnicamente, se você entra em um site qualquer que te apresenta um certificado e nenhuma mensagem é mostrada para você, é porque não há problemas (aparentes) com o certificado enviado para você.

### 10.8.3. PKI – Public Key Infrastructure – infraestrutura de chaves públicas

A alma da certificação digital é a confiança. Confiar em um certificado requer confiança na autoridade que o emitiu. Por sua vez, confiar na autoridade que o emitiu requer que se confie na autoridade que emitiu o certificado para aquela autoridade e assim por diante.

Não encontramos apenas AC + Certificado, não! Às vezes, é necessário confiar em vários níveis para que o certificado seja válido.

O ambiente, formado por vários níveis (ACs, Usuários, ACs Raiz, ARs) é chamado **PKI**, ou **ICP (Infraestrutura de Chaves Públicas)**. Uma PKI é um conjunto de regras, técnicas, práticas e procedimentos que existe para gerar garantias aos seus usuários. Uma PKI é, em suma, uma grande “cadeia de confiança” em que, quando se confia em um de seus componentes, se está confiando também em toda a PKI.

Uma PKI é formada, normalmente por:

- **Uma Autoridade Certificadora Raiz (AC Raiz):** o topo da PKI, todos os componentes da PKI confiam na AC Raiz. Ela é o gênese de toda a confiança. A AC Raiz emite certificados atestando a autenticidade das AC intermediárias e a AC Raiz é autocertificada (ou seja, ela emite seu próprio certificado), mas a AC Raiz não pode emitir certificados para usuários finais (como sites ou pessoas físicas).
- **Autoridades Certificadoras Intermediárias (AC, simplesmente):** são subordinadas à AC Raiz e têm seus certificados emitidos por esta. As ACs são entidades públicas ou privadas com estrutura física segura o suficiente para guardar, sigilosamente, os dados de seus clientes (certificados).

É justamente a AC intermediária que emite os certificados para os usuários finais (como os certificados para e-mail que usamos ou os certificados que autenticam os sites que acessamos).

- **Autoridades de Registro (AR):** é uma instituição, associada a uma AC, que recebe as solicitações de emissão de certificados dos usuários que os requerem. Uma AR é o “posto de atendimento” da PKI, digamos assim.

A AR não pode emitir certificados, mas pode receber o solicitante, cadastrar sua requisição, receber os dados documentais da pessoa/empresa/site, e, garantindo que a solicitação é válida (por exemplo, verificando os documentos e a pessoa que solicita o pedido de emissão do certificado), enviar o pedido de emissão do certificado para a AC.

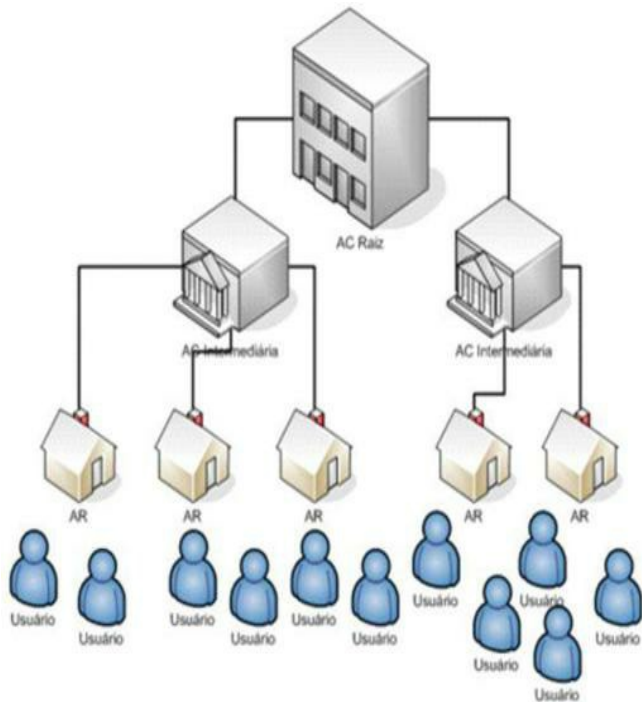
Todos os pedidos relacionados ao certificado (emissão, revogação, renovação) são feitos pelo titular do certificado à AR. A AR analisa o pedido e, confirmando sua autenticidade, repassa a solicitação à AC responsável.

- **Usuários:** pessoas/empresas/sites que se beneficiam do “universo de confiança” da PKI, são os componentes que solicitarão e utilizarão certificados emitidos pelas ACs daquela PKI. Confiar em um certificado de um usuário qualquer requer confiança em todos os níveis daquele certificado, desde a AC Raiz.

- **DPC (Declaração de Práticas de Certificação):** documento, ligado às AC, que especifica os critérios técnicos da estrutura de certificação, como: algoritmos usados, tamanho das chaves, validade dos certificados, tempo máximo de publicação da LCR (Lista dos Certificados Revogados), critérios para a emissão dos certificados, documentos exigidos etc.

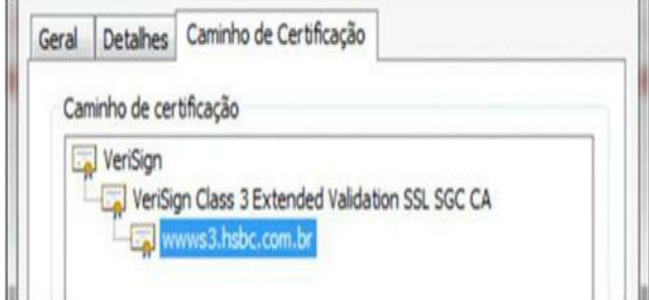
A DPC é revisada e republicada várias vezes, e suas diretrizes determinam como ela vai agir para emitir, revogar ou renovar certificados – se duas partes em um contrato concordam com a DPC da autoridade certificadora, então será fácil confiar na AC e em todos os seus certificados.

Veja, a seguir, um esquema da organização de uma PKI (e seus componentes):



*Figura 10.14 – Estrutura de uma PKI.*

Dê uma olhada na árvore de certificação do certificado do Banco HSBC (emitido para garantir a autenticidade do site do banco).



*Figura 10.15 – Detalhe do certificado do HSBC.*

#### **10.8.4. E a certificação digital do ponto de vista jurídico?**

Diversos países têm suas ideias a respeito da certificação digital. O Brasil, por incrível que pareça, é um dos mais “avançados” nesse sentido. Desde 2002, com a aprovação da Medida Provisória 2200, os processos certificadores digitais são aceitos “plenamente” em comunicações com o Governo pela Internet.

Claro que implantar isso é realmente um processo gradativo e demorado, tanto que mesmo hoje, onze anos depois, ainda não é toda instituição do Governo Federal, ou das outras escalas, que aceita os certificados digitais. (Não por se negarem, mas por não implementarem recursos de informática que possibilitem isso.)

Quando uma transação eletrônica é feita com um certificado, tem-se a integridade da realização da transação, a identidade do usuário (autenticidade) e isso com o aval de uma instituição confiável (a AC).

Por essa razão, tem-se a garantia do não repúdio (non-repudiation). Ou seja, o autor da transação não pode dizer que não foi ele quem a realizou ou que ela não foi realizada. Essa é a base para a validade jurídica da transação eletrônica.

Mas, as transações eletrônicas com os diversos órgãos do Governo Federal só são consideradas válidas (juridicamente, inclusive) quando os usuários utilizam certificados válidos para a PKI do Governo Federal, chamada **ICP-Brasil**, instituída na Medida Provisória citada anteriormente.

#### **10.8.5. A ICP-Brasil**

A ICP-Brasil, ou Infraestrutura de Chaves Públicas Brasileira, é a “rede de confiança” (PKI) adotada pelos órgãos do Governo Federal. A ICP-Brasil é gerenciada pelo Instituto Nacional de Tecnologia da Informação – ITI, que é uma autarquia federal vinculada à Casa Civil da

Presidência da República. O ITI é a AC Raiz da ICP-Brasil (também conhecida como Autoridade Certificadora Raiz Brasileira).

Além da AC Raiz da ICP-Brasil, que é o ITI, existem diversas ACs intermediárias, para quem a AC Raiz emite certificados, demonstrando confiança neles: Caixa Econômica, Receita Federal, Serasa, Serpro, Certisign etc. Por sua vez, as AC Intermediárias podem emitir certificados para pessoas físicas, empresas ou outras AC (pode haver outros níveis de AC intermediárias, tantos quanto a PKI confiar – ou determinar, em sua DPC).

As normas de funcionamento da ICP-Brasil são determinadas pelo Comitê Gestor ICP-Brasil e repassadas a todas as ACs e ARs. A AC Raiz emite certificados para as AC intermediárias sob a autorização do Comitê Gestor, que realiza a auditoria nas candidatas a AC para autorizar a emissão dos certificados delas.

Para se tornar uma AC registrada na ICP-Brasil, é necessário passar por uma análise muito rígida, determinada pelo Comitê, que analisa todos os aspectos da instituição a ser registrada, aprovando-a somente se passar pelos mais diversos critérios.

#### **10.8.6. Como emitir um certificado?**

O processo de emissão do certificado digital segue alguns passos:

1. O usuário que deseja possuir um certificado deve comparecer a uma AR pessoalmente munido dos documentos necessários que comprovem sua identidade: RG, CPF, comprovante de residência e quaisquer outros documentos que a AR julgar necessários (isso está estabelecido na norma da AC Raiz e das AC Intermediárias).
2. A AR cria, através de software e algoritmo específicos, o par de chaves daquele usuário (a chave pública e a chave privada), usadas para assinatura digital ou para a criptografia das mensagens de e-mail.
3. A AR envia a chave pública do usuário para a AC, que, assina a chave pública do usuário e seus dados com a sua chave privada (da AC) e devolve o certificado pronto à AR.
4. A AR armazena o par de chaves do usuário no computador dele (ou envia um e-mail a este com estas chaves). Há também a possibilidade de colocar o par de chaves em um “hardware criptográfico” que é uma memória com esse certificado.
5. A partir daí, o usuário poderá sempre usar aquela chave privada para assinar as mensagens de e-mail que ele criar para outros. A mensagem assinada conterá o seu certificado (e o certificado conterá a chave pública dele).

Um certificado é um documento (um conjunto de bytes) e, por isso, pode ser armazenado em diversas mídias (como no HD do computador do usuário, num cartão especial ou numa memória conectada à porta USB – chamada token criptográfico). Conheça o cartão (smart card) e o token USB.

A figura abaixo demonstra smart cards (cartões) e token USB (parecido com um pen drive) emitidos pela Certisign (uma das mais confiáveis, se não a mais confiável, AC do país).



Figura 10.16 – Smart cards criptográficos (armazenando o e-CPF) e token USB.

10.8.7. Exemplos de uso de certificados ICP-Brasil Através da utilização de certificados digitais, é possível assinar e-mails, ser reconhecido por sistemas de empresas e até mesmo por sites.

Atualmente, o governo brasileiro faz uso de diversos certificados digitais emitidos, claro, sob as normas da ICP-Brasil:

- **e-CPF:** certificado para contribuintes pessoas físicas terem acesso a seus dados e realizar processos de declaração e retificação de declaração de Imposto de Renda, entre outros, junto à Receita Federal do Brasil.
- **e-CNPJ:** certificado digital para empresas – normalmente usado pelos contadores destas empresas, para comunicação eletrônica junto à Receita Federal do Brasil.
- **NF-e:** certificado digital usado por empresas para a emissão de notas fiscais eletrônicas.

## 10.9. Então, em resumo...

A comunicação pela Internet é basicamente anônima e insegura. Não se tem certeza de quem

está do outro lado da rede nem de que as mensagens recebidas realmente foram escritas daquele jeito.

Os processos matemáticos de criptografia e hash dão, aos usuários e empresas, garantias técnicas a respeito da autenticidade, sigilo e integridade desses dados. E foram justamente as tecnologias de criptografia assimétrica (aquela que usa um par de chaves para cada usuário) e de hash que deram origem aos processos de assinatura digital, associando, de forma unívoca, um par de chaves a um usuário/empresa.

A certificação digital vem somente adicionar mais seriedade e confiabilidade nesse sistema, atestando, por meio das AC, que a chave pública de Fulano realmente pertence a Fulano (isso é feito através do certificado, que contém a chave pública de Fulano e é assinada pela AC, garantindo isso).

A certificação digital é a forma mais segura, hoje em dia, de permitir que as transações realizadas pela Internet sejam seguras e garantidas, evitando, assim, que, para entrar em contato com o Governo, ou outra instituição, seja necessário dirigir-se a ele e, em muitos casos, enfrentar filas enormes e todos aqueles inconvenientes inerentes ao serviço “corporal”. Através do uso de certificados, futuramente, não precisaremos sair de casa para fazer absolutamente nada em relação ao governo e, mais futuramente ainda, às empresas com quem costumamos lidar fisicamente.

## 10.10. Q u e s t õ e s – S e g u r a n ç a

1. O golpe de Pharming é caracterizado por:
  - a) o envio de pacotes TCP/IP de tamanho inválidos para servidores, levando-os ao travamento ou ao impedimento de trabalho;
  - b) a impossibilidade de identificação do número de IP de máquina conectada à rede. Dessa forma, muitos dos serviços de segurança disponíveis deixam de funcionar, incluindo os “rastreamentos” que permitem a identificação de segurança das fontes de origem de ataques;
  - c) instalar em um computador conectado a uma rede um programa cliente que permite a um programa servidor utilizar esta máquina sem restrições;
  - d) a alteração dos registros do servidor de nomes, direcionando-os para endereços de páginas fraudulentas;
  - e) a captura de “quadros” em redes de difusão (como a Ethernet).
2. Acerca dos certificados digitais, assinale a alternativa incorreta.
  - a) Permitem identificar usuários que enviam e-mails.
  - b) São emitidos e renovados por Autoridades de Certificação (AC).
  - c) Contêm as chaves públicas e privadas de seus titulares.
  - d) Podem ser revogados.
  - e) Apresentam data de validade.
3. A técnica de MAC spoofing é especialmente interessante para:
  - a) possibilitar o sniffing entre redes distintas;
  - b) impedir que se descubra o endereço IP do atacante na Internet;
  - c) possibilitar o sniffing em redes que usam Switches;
  - d) permitir o ataque de SYN flooding em redes locais que não usam TCP/IP;
  - e) proteger uma conta de e-mail do SPAM.
4. Um filtro de pacotes não é capaz de:
  - a) detectar pacotes inválidos por seu endereço IP de origem;
  - b) bloquear pacotes oriundos de um determinado servidor na Internet;
  - c) permitir a passagem de pacotes de um determinado serviço, como a Web;
  - d) identificar comportamentos condizentes com SYN Flooding e Phishings;
  - e) bloquear a entrada e a saída de pacotes de FTP.
5. Considere que em uma comunicação entre um usuário A e um usuário B há um determinado usuário ilegítimo C que consegue interceptar as mensagens entre ambos e reenviá-las como se fosse o legítimo interlocutor. Dessa forma, A pensa que C é B e B pensa que C é A. Essa técnica de espionagem é conhecida como:
  - a) ARP poisoning;
  - b) DNS poisoning;
  - c) E-mail poisoning;



- d) Man-in-the-middle;
- e) Ping of Death.

6. (Adaptada) Em um sistema de criptografia simétrica:

- a) todos os usuários usam chaves diferentes;
- b) o processo de criptografia demora mais tempo que num sistema assimétrico;
- c) é usada uma mesma chave para encriptar e decriptar mensagens;
- d) é usada uma chave para encriptar e outra para decriptar mensagens;
- e) não são usadas chaves nesse processo de encriptação, tornando o sistema muito mais rápido.

7. (Adaptada) A presença do ícone do cadeado na barra de status de um navegador indica que:

- a) o site está se comunicando de forma criptografada com o seu navegador;
- b) o site não tem certificados digitais;
- c) o site está localizado na zona de sites proibidos no seu navegador;
- d) o site está fechado;
- e) o firewall está bloqueando o conteúdo daquele site.

### **11.1. Considerações iniciais**

Bom, caro leitor.

Começamos, aqui, mais um capítulo do livro. (Este, em minha opinião, o mais simples de todos!) Gostaria de avisar, apenas que esse assunto não é exigido em grande quantidade na maioria das bancas examinadoras. Backup é quase que uma exclusividade da Esaf!

Mesmo assim, não custa muito dar uma olhadinha neste capítulo se o conteúdo programático do seu edital mencionar algo como “Noções de Backup (Cópia de Segurança)”, embora a presença desse tópico no edital não garanta a presença de questões desse assunto na prova. Aproveite.

### **11.2. Noções básicas sobre backup**

Backup (lê-se “becáp”) é a operação de copiar arquivos por segurança em um local (disco ou memória) diferente do original.

Se um determinado usuário possui arquivos importantes em seu computador, ele poderá copiá-los para um disquete ou um CD para que, caso haja algum problema com as informações originais (falha no micro, por exemplo), ele possa recuperá-las posteriormente (porque as tem gravadas em outro lugar).

Por mais banal que pareça, o simples ato de copiar um arquivo do seu disco rígido para um pen drive, e guardá-lo por segurança (note o intuito de tê-lo armazenado em outro lugar POR SEGURANÇA), já constitui um processo de backup. Copiar os arquivos de um computador para um CD ou DVD virgem (para guardar tais discos) também é um backup.

### **11.3. Backups em concursos**

Em primeiro lugar, em se tratando de provas de concursos para quem não é da área de informática, só é comum encontrar alguma menção mais detalhada sobre backups nas provas da ESAF. Outras bancas examinadoras simplesmente esquecem-se desse assunto, mesmo descrevendo sua existência nos conteúdos programáticos dos editais.

Ainda tem mais! Para as provas em que assunto de backup é exigido, consideram-se todos os detalhes de funcionamento da “forma oficial” de fazer backup. Em suma, é necessário aprender como o backup é feito nas empresas, pelos profissionais treinados, e não como nós fazemos em casa – que é o que chamo de “backup artesanal” (quando a gente lembra de fazer, né? O que é MUITO RARO!)

#### **11.3.1. Para que o processo de backup é usado?**

Fazemos backups para manter cópias seguras de nossos dados. (As bancas examinadoras, como Esaf, FCC, Cespe etc. sabem disso!)

Em resumo, backups são feitos para que possamos ter nossos documentos (textos, fotos, planilhas, desenhos, músicas) em local seguro (NECESSARIAMENTE DIFERENTE do disco de

origem) porque se houver algum problema com os dados originais em nosso disco, será possível recuperá-los.

Preste bem atenção a isso: qualquer questão de prova que diga que o backup tem que (ou pode) ser feito NO MESMO DISCO/MÍDIA onde os dados estão atualmente, isso é FALSO! Backups devem ser feitos, impreterivelmente, em discos/mídias/servidores diferentes daqueles onde o dado está atualmente armazenado.

### **11.3.2. E para que o processo de backup não é usado?**

Existem muitas questões engraçadas sobre isso. Em geral, backups não são feitos para guardar cópias dos programas. Há questões que sugerem que os backups são usados (ou deveriam) para guardar cópias do sistema operacional, ou dos programas de escritório, ou ainda do BIOS da placa-mãe (é, eu sei, essa é até difícil para “engolir”).

Em primeiro lugar, todos os programas que temos em nossos computadores são reinstaláveis. Basta o usuário ter o DVD (ou pen drive hoje em dia) do programa que deseja reinstalar na máquina e ele será reinstalado. Não há necessidade de fazer backups do Windows, nem do Word, nem de nenhum outro programa.

Em segundo lugar, se alguém vier falar de “backup do BIOS”, está “viajando”. Lembre-se de que o BIOS é um firmware (programa gravado num chip firme, como de memória ROM) – ele está lá para nunca sair, nunca ser apagado. Portanto, nem se prevê a mínima possibilidade de que o BIOS precise de backup (é simplesmente absurdo!).

Qualquer questão que mencione isso (ou seja, backups de programas) é FALSA! Backups são feitos para os DADOS (nossas informações).

## **11.4. Conhecendo o processo de backup**

Basicamente, para entender todos os processos relacionados a backup, é preciso saber duas coisas:

### **1. Backup é um processo cíclico.**

Isso significa que, nas empresas, um backup é visto como algo que tem começo, meio e fim (e novo começo depois). Ou seja, há planejamento, definições e regras a serem seguidas para que, inclusive, de vez em quando, seja reiniciada toda a rotina de backup na empresa.

Um novo ciclo de backup sempre reinicia com um Backup Total (também chamado de Backup Normal).

### **2. O backup tem de se manter contemporâneo aos dados que se deseja assegurar.**

Em relação à característica 1, essa aqui tem sentido de: “somente o último ciclo de backup (o mais recente) é importante para a recuperação dos dados em caso de problemas”. É bom lembrar, caro leitor, que o intuito do backup não é guardar dados por 10 ou 20 anos. O objetivo do backup é manter cópias atualizadas dos arquivos do sistema. Só são importantes os dados mais recentes!

Backups feitos há dois ou três anos não são necessários (se levarmos em consideração uma técnica de backup que usa, inteligentemente, esses ciclos de backup).

## **11.5. Onde os backups são feitos?**

Ainda no intuito de conhecer como os backups são feitos em empresas, profissionalmente, que é o que realmente importa para a ESAF (e as demais bancas, provavelmente), é necessário que se saiba que as cópias dos dados da empresa são feitas, normalmente, em fitas magnéticas (as famosas “fitas DAT”).

Lembrando, como está descrito no capítulo sobre hardware, que DAT é apenas um dos tipos de fitas magnéticas usadas. E um tipo muito antigo! Hoje, o termo mais comum para fazer referência a esse tipo de mídia é “fitas magnéticas” (ou “fitas para backup”).



*Figura 1.11 – Uma Fita DDS (da HP) – um dos formatos mais usados atualmente.*

Essas fitas, como já deu para perceber, são magnéticas e, por isso, são reutilizáveis. Sim! Assim como qualquer outra mídia magnética, as fitas são regraváveis. As fitas podem ser utilizadas novamente sem problema se o ciclo dos backups for bem determinado.

Aí voltamos novamente ao início de tudo: o ciclo! Veremos, com o passar desse assunto, que essa história de ciclos é mais importante que parece.

Então, atente para uma coisa simples sobre as fitas: se uma delas for usada hoje para guardar o backup dos dados que foram alterados hoje na empresa e depois essa mesma fita for usada daqui a 10 dias, os dados novos (os dados que aparecerão daqui a 10 dias) substituirão completamente os dados de hoje.

Sim! A fita não pode ser usada “por partes”. Quando uma fita é usada, desconsidera-se tudo o que ela trazia antes. Você deve entender que quando uma fita é reutilizada, os dados que ela

continha são completamente apagados! Portanto, só se pode reutilizar uma fita se houver a certeza de que dados lá contidos já não são mais necessários.

Vamos ver, durante o estudo dos tipos de backup e de suas principais estratégias, que é fácil saber quando os dados de uma fita não são mais necessários.

### **11.6. Como os backups são feitos?**

Basicamente, há duas etapas antes de iniciar o processo de cópia em si:

1. Escolher quais são os arquivos (e pastas) que deverão ser backupeados.
2. Escolher o tipo do backup a ser realizado. (Isso determinará quais dentre aqueles arquivos escolhidos serão realmente colocados na fita.)

Depois disso, basta iniciar o processo de backup, deixando a cargo do programa de backup a escolha dos arquivos que realmente serão copiados.

“Quer dizer que mesmo que eu escolha uma pasta com 1.000 arquivos, pode ser que nem todos eles sejam colocados na fita?” (Você pode estar se perguntando agora.)

Sim! Pode ser que dos 1.000 arquivos, apenas 50 efetivamente sejam copiados para a fita. Essa escolha é feita por uma série de critérios, que diferem entre si de acordo com o tipo de backup que foi escolhido.

“João, e se eu simplesmente copiar os arquivos? Ou seja, se arrastar os arquivos normalmente, usando o Windows Explorer?”

Caro leitor, você está se referindo ao que chamei, há pouco, de “backup artesanal”? Bom, se você o fizer, todos os arquivos serão copiados, mas é bom que se entenda: esse “processo” não é O BACKUP que é considerado para os concursos.

Para os concursos, subentende-se que está sendo usado um programa específico para fazer backups. Esses programas oferecem várias opções de procedimentos de backup, bem como realizam certas “operações” que nossas cópias amadoras não fazem (o ato de “marcar” os arquivos que foram copiados, como veremos).

### **11.7. Programas para backup**

Há vários programas que ajudam o usuário a fazer backups. Esses programas normalmente acompanham os próprios sistemas operacionais instalados no computador. Vamos usar como exemplo nas fotos a seguir o programa Microsoft Backup e Restore Center, que acompanha o sistema operacional Windows® 7 (exceto versão STARTER EDITION).



## Configurar backup

Você deseja fazer backup de quais itens?

Marcar a caixa de seleção dos itens a serem incluídos no backup. [Que arquivos são excluídos do backup por padrão?](#)

- ▲ Arquivos de Dados
  - ☒ Fazer backup de dados para usuários recentemente criados
  - ▶ ☒ Bibliotecas de Joao Antonio
- ▲ Computador
  - ▶ ☐ Principal (C:)
  - ▶ ☐ RECOVERY (F:)

☐ Incluir uma imagem do sistema das unidades: RECOVERY (F:), Principal (C:)

O local de backup selecionado não é compatível com a criação de imagens de sistema.

Avançar

Cancelar

Figura 11.2 – Microsoft Backup e Restore Center, do Windows 7.

## 11.8. A benta “marcação” dos arquivos – teoria original de backup

Se tem uma coisa que realmente incomoda os concursandos em geral é, entender, de forma simples, as confusões relacionadas à “marcação” dos arquivos que passaram por um backup. Vou tentar minimizar essa estranheza agora, caro leitor.

Pense comigo, amigo leitor, imagine-se como Técnico ou Auditor da Receita Federal e que tenha sido escalado para “conferir” os conteúdos de 3.000 caixas em um contêiner qualquer de um navio. Você acha que conseguirá conferir todas as 3.000 caixas em um dia só?

Bom, a menos que quem esteja lendo este parágrafo agora seja o “The Flash”, acho que o serviço vai render por alguns dias, não é? É justamente aí que está... Como saber, na terça-feira, quais as caixas que já foram conferidas na segunda-feira?

Simples! Depois de conferir qualquer uma das caixas, é só “marcar” a caixa, não é mesmo? Essa marcação pode ser de várias formas: fazendo um “X” na caixa; arrancando algum pedaço de sua tampa; arrancando algum adesivo; colocando algum adesivo. É você, conferente, que decide qual FORMA será usada para indicar as caixas já conferidas.

Portanto, caro leitor, quero que entenda que a “marcação” nada mais é que uma forma de o programa que faz o backup (e no nosso exemplo seria o “conferente”) saber quais arquivos já passaram por um backup e quais aqueles que precisam ser copiados no próximo backup.

Então é necessário, para todos os fins do estudo teórico do backup, saber que:

- Quando um arquivo é criado, ele não tem a marca. Isso significa que todo arquivo, quando nasce, já nasce “precisando” passar por um backup.
- Quando um arquivo é copiado em um processo de backup, ele é marcado (isso dependerá, também, do tipo de backup usado). Ou seja, quando um arquivo é backupeado, ele recebe uma indicação, legível pelo programa de backup, que diz que ele acabou de passar por um backup e não tem que passar, necessariamente, por outro.
- Quando um arquivo é alterado (modificado e salvo), ele perde a marcação. Sim, isso é lógico porque a cópia que foi alterada no micro do usuário não é mais igual àquela cópia que está no backup. Logo, o arquivo alterado precisará, sim, passar pelo próximo ato de backup.

Aqui estão alguns arquivos antes e depois de um backup normal (um dos tipos de backups que “marca” os arquivos):



*Figura 11.3 – Arquivos antes de um backup normal.*



*Figura 11.4 – Os mesmos arquivos depois de um backup normal.*

É desnecessário dizer, também, que a marcação não aparece tão “descaradamente” como se



vê nessa figura. Nos sistemas operacionais da família Windows, essa marcação é registrada em um atributo com que normalmente não nos preocupamos, chamado “Atributo de Arquivamento” ou “Atributo de Arquivo”.

E esse “Atributo de Arquivamento” é outro palco de confusões, veremos o porquê.

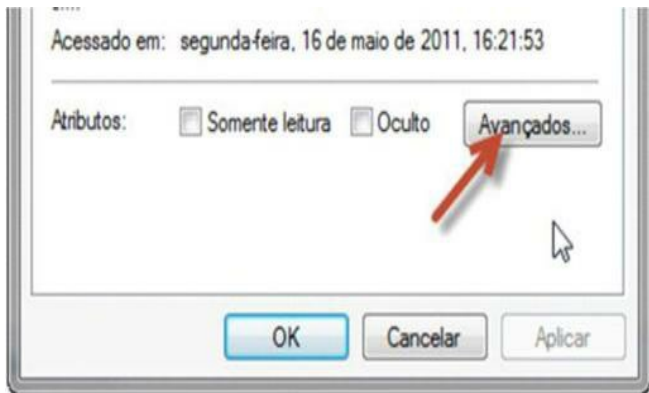
### 11.9. Atributo de arquivamento – a “marcação” do Windows

Vimos que é muito lógico haver algum procedimento de “marcação” dos arquivos que foram copiados em um processo de backup. Veremos mais adiante que nem todos os processos de backup marcam os arquivos (isso é “privilegio” apenas de alguns tipos).

Vimos também, e vamos explorar mais a fundo, que, no Windows, os arquivos são indicados como tendo passado por um backup através de um atributo, uma informação, chamado “Atributo de Arquivamento”. Mas onde ele está? O que ele é exatamente?

Para encontrá-lo, basta clicar com o botão direito do mouse em qualquer arquivo no Windows Explorer e selecionar a opção “Propriedades” (ou, depois de selecionar o arquivo, clicar no menu Arquivo/Propriedades).

Na janela de propriedades, será possível encontrar, na parte inferior, o botão AVANÇADOS, como visto a seguir, para dar acesso aos atributos mais técnicos do arquivo em questão:



*Figura 11.5 – Botão AVANÇADOS – encontrado na janela de Propriedades do Arquivo.*

O Atributo de Arquivamento (ou Arquivo, ou Arquivo-Morto) é encontrado justamente dentro da janela aberta a partir do botão AVANÇADOS... conforme mostrado na figura a seguir.

## Atributos Avançados



Escolha as configurações para esta pasta.

### Atributos de arquivos

- ☒ O arquivo está pronto para ser arquivado
- ☒ Permitir que o conteúdo deste arquivo seja indexado junto com as propriedades do arquivo

Figura 11.6 – O Atributo de Arquivamento, finalmente.

A caixa de verificação “O arquivo está pronto para ser arquivado” é o que nós conhecemos mais intimamente como “Atributo de Arquivamento”. É essa informação que diz se um arquivo está apto a passar pelo próximo backup ou se já passou.

Ou seja, os programas de backup dos sistemas Windows costumam “ler” esse atributo durante a realização de um backup, para saber se os arquivos precisam, ou não, passar por aquele processo. E é aí que “mora o perigo” para quem estuda!

O atributo de arquivamento é meio “ao contrário” da ideia de “marcação” a que fomos apresentados há pouco na teoria “oficial” do backup. Ou seja, é ele que indica a relação do arquivo com os programas e hábitos de backup, mas ele faz isso “de forma invertida”. Vamos entender:

- a. Quando um arquivo é criado, ele já nasce com o atributo de arquivamento marcado (aquela caixinha “o arquivo está pronto para ser arquivado” já nasce marcada).
- b. Quando um arquivo é backupeado (ou seja, quando ele passa por um backup), seu atributo de arquivamento é desmarcado.
- c. Quando um arquivo é alterado, o atributo de arquivamento volta a ser marcado.

Então, você percebe que “o arquivo é marcado como tendo passado por um backup” é a mesma coisa que dizer “o atributo de arquivamento do arquivo foi desmarcado”. É um pouco estranho, mas note uma coisa. A expressão “o arquivo está pronto para ser arquivado” já deixa tudo claro, não é? Ela quer dizer, em poucas palavras, que o arquivo deve passar pelo próximo backup (se a caixa estiver marcada).

**Aí você pergunta: “E então, João, como devo interpretar na prova? Os arquivos são**

**marcados ou desmarcados?” – É bem simples! Vamos às explicações baseadas no que já se viu em questões de provas de concursos.**

Quando o enunciado da questão citar “... o arquivo é marcado...”, devemos entender que o elaborador da questão faz referência à teoria original de backup que vimos (em que “marcado” significa “ter passado por um backup”). Porém, quando vemos expressões como “... o atributo de arquivamento (ou ‘de arquivo’) é marcado...”, levamos em consideração que o elaborador da questão está se referindo ao “jeito Windows” de entender backup.

Ou seja, quando estiver mencionando “... o arquivo será marcado...”, lemos tudo como na teoria original de backup:

1. Arquivo marcado é sinônimo de “arquivo que já passou por um backup” (e que não precisa passar mais).
2. Arquivo desmarcado significa “arquivo criado ou alterado recentemente” (e que tem de ser copiado no próximo backup).

E, quando a expressão na questão contiver algo como “... o atributo de arquivamento é desmarcado (ou marcado)...”, ou seja, quando falarem no atributo, devemos entender que o examinador vai ler a questão conforme o entendimento do Windows:

1. Atributo marcado significa “arquivo que foi criado ou alterado recentemente” (e que, por isso, tem de ser copiado no próximo backup).
2. Atributo desmarcado é o mesmo que “arquivo que acabou de passar por um backup” (e não precisa mais passar por outro).

Fica fácil entender que o “ponto de vista” da questão está relacionado à forma da pergunta. Mencionando o “atributo de arquivamento”, entende-se toda a questão como descrita com base no uso do atributo. Em não mencionar o atributo, o elaborador deixa claro que quer que se cite a ideia de “marcado” e “desmarcado” com base na teoria inicial de backup.

Algumas questões recentes, para evitar interpretações dúbias, deixaram claras suas posições quando mencionaram, em seus enunciados, explicações nas duas “ideias”, como veremos a seguir num texto de uma questão recente da ESAF:

“Quando um Backup Normal é realizado, todos os arquivos são copiados e depois disso são marcados, ou seja, eles têm seus atributos de arquivo desmarcados.”

Notaram? O elaborador da questão citou os dois métodos. Não tem como deixar dúvidas. Agora, falando em “backup normal”, vamos aos tipos de backups e suas principais características.

## **11.10. Tipos de backups**

Há alguns tipos comuns de processos de backup, cada um com suas características bem definidas em relação a outros.

Apenas como um lembrete: é bom citar que os backups serão sempre realizados em fitas magnéticas, ok? Portanto, se eu mencionar, daqui em diante, expressões como “... foi copiado para a fita...”, você não vai estranhar, não é, leitor? Vamos lá!

### **11.10.1. Backup normal (ou global)**

Quando o usuário seleciona uma pasta (ou várias pastas) e aciona o backup normal, todos os

arquivos selecionados (ou seja, contidos nas pastas selecionadas) serão copiados para a fita. Sem distinção! Ou seja, “todo gato é pardo”!

Mesmo que haja arquivos que já passaram recentemente por backups, ou seja, que estejam marcados (que, no Windows significa ter o “atributo” desmarcado), o backup normal copia todos. Serão copiados para a fita todos os arquivos selecionados, tendo, ou não tendo, a marcação.

Depois que o backup normal é realizado, todos os arquivos copiados são marcados (ou seja, têm seus “atributos de arquivamento” desmarcados), como prova de que acabaram de passar por um backup.

Dica: este aqui é o Backup principal! Sempre, em uma estratégia definida para uma empresa, a rotina de backups começa por um exemplar deste aqui! Os “ciclos” (ainda vamos conhecer) são sempre iniciados por um backup normal.

### **11.10.2. Backup incremental**

Este tipo de backup complementa o backup normal. Quando selecionamos algumas pastas para fazer backup e escolhemos, no programa de backup, a realização de um incremental, nem todos os arquivos selecionados serão copiados.

Apenas serão copiados aqueles arquivos criados ou alterados desde o último backup (ou seja, os arquivos que “precisam ser copiados”). Isso, claro, se refere aos arquivos que não têm a marcação, ou, em se tratando de Windows, aqueles arquivos que têm o atributo de arquivamento marcado.

Portanto, se em uma pasta qualquer a ser backupeada existirem 1.000 arquivos, dos quais 300 não foram alterados desde o último backup, na realização de um backup incremental, somente os 700 arquivos desmarcados (ou seja, com o atributo marcado) serão copiados para a fita.

NOTE: não é você quem indica os 700 arquivos. Você indica a pasta inteira (que contém os 1.000 arquivos). O próprio programa de backup selecionará os arquivos que “merecem” ser backupeados, de acordo com o fato de indicarem isso (com o “atributo”).

Depois de copiar os arquivos para a fita, seus originais são marcados (ou seja, têm o atributo de arquivamento desmarcado), indicando que eles acabaram de passar por um processo de backup.

### **11.10.3. Backup diferencial**

O backup diferencial é semelhante ao backup incremental no que se refere a “quem será copiado”. Ou seja, esse backup também copia apenas os arquivos que precisam ser copiados (aqueles que foram criados ou alterados desde o último backup).

Em outras palavras, no diferencial, só são copiados os arquivos que não estão marcados (ou que, segundo o Windows, têm o atributo de arquivamento marcado).

Como é muito parecido com o incremental, costumamos dizer que “um faz o que o outro faz” e, por isso, nas estratégias que vemos funcionando nas empresas por aí, eles são excludentes: quem usa incremental, não usa diferencial, e vice-versa!

O diferencial, porém, faz outra coisa com os arquivos depois que os copia: nada! Sim! Por mais estúpido que pareça, o backup diferencial não marca os arquivos (ou seja, não desmarca o atributo de arquivamento). O diferencial simplesmente copia os arquivos e não faz nenhum tipo

de indicação que os copiou.

Calma. Você já deve estar imaginando o “auditor” que confere as caixas, mas não faz nenhum tipo de “marca” nelas. Quando esse profissional chegar no dia seguinte ao contêiner para continuar a conferência das caixas, como ele procederá?

Não quero nem saber. Vamos ver para que esse backup serve mais adiante.

#### **11.10.4. Backup diário**

Através da realização de um backup diário, são copiados para a fita apenas os arquivos que foram criados ou alterados numa data específica (normalmente naquela data em que o backup está sendo realizado). Ou seja, no momento da realização do backup, o programa pergunta ao usuário qual a data em que ele quer que o backup seja feito.

É possível fazer backups diários retroativos? Sim! Mas não é uma coisa muito inteligente. Imagine-se estando, digamos, no dia 10 de agosto e você quer fazer o backup do dia 7 de agosto (o dia mais bonito do ano!) porque se esqueceu de fazê-lo no dia certo. O que acontecerá?

Bom, todo arquivo carrega consigo sua “data de última modificação” (e só ela!). Um arquivo não guarda sua “data de penúltima (ou antepenúltima) modificação”. Diante disso, o backup diário só será capaz de analisar o “merecimento” de um arquivo por sua última data.

E se um dos arquivos que foi modificado no dia 7 de agosto também foi modificado posteriormente em 8 ou 9 de agosto? Ao fazer, no dia 10, um backup diário dos arquivos modificados no dia 7, esse arquivo não será contemplado. Se o backup fosse feito no dia 7 (ou, no mais tardar, nas primeiras horas do dia 8), esse arquivo seria perfeitamente copiado na fita.

Depois que um backup diário é realizado, ele não marca nenhum dos arquivos. Ou seja, ele nem se importa com quem “tem marcação” ou “não tem marcação” e nem faz nenhum tipo de alteração nesse sentido! Todos os arquivos copiados para a fita permanecerão como estavam antes de o backup ser realizado.

#### **11.10.5. Backup de cópia**

O backup de cópia não é citado em nenhuma estratégia ou rotina de backups. As empresas não “preveem” o uso desse tipo de backup! Ele é considerado um “backup emergencial”, feito para momentos em que se precisa fazer um backup completo que não altere a rotina da empresa.

Exemplo: em uma empresa, vai-se instalar um novo programa no servidor de e-mails. Esse programa novo promete facilitar e melhorar a vida dos usuários de e-mail da empresa. Antes, porém, de instalar qualquer coisa, deve-se, por precaução, fazer um backup de cópia dos dados existentes naquele servidor.

Já imaginou se o backup não é feito e a instalação do “novo programa” que promete “melhorar o uso do e-mail” acaba detonando o servidor? Já imaginou se o novo programa termina de lascrar tudo lá dentro? E os e-mails dos usuários? E os documentos? E as configurações?

Você pode até perguntar: “João, por que não fazer um backup normal, que copia tudo? O backup de cópia é igual ao normal?” – Não! O backup de cópia tem uma pequena diferença que o torna ideal para momentos emergenciais e não previstos como esses.

O backup de cópia, quando realizado, copia todos os arquivos, sem distinção (nesse quesito, é

idêntico ao normal). Mas, depois de feita a cópia, esse backup NÃO marca nem desmarca ninguém (deixa todos como estão). “E por que isso, João?”

Simples! Por ser emergencial, ele tem de copiar todo mundo mesmo! E mais, ainda por ser emergencial (e, portanto, não previsto na rotina da empresa), ele não pode marcar ninguém, para não “alterar” a rotina de backup utilizada naquela empresa.

Caso haja problemas na instalação do “novo programa”, o backup de cópia está lá, para recuperar os dados mais recentes. Todavia, se não houver nenhum problema com a instalação do novo programa, o backup de cópia não terá alterado nenhum dos arquivos, não influenciando nos resultados da rotina de backup da empresa.

### 11.11. Resumo dos tipos de backup

Para que se possa memorizar facilmente todos os tipos de backup, basicamente é necessário saber quem será copiado pelo backup e o que o backup faz com os arquivos depois de copiados. Isso está compilado na tabela a seguir.

| O backup... | Copia...                       |  |
|-------------|--------------------------------|--|
| Normal      | Todos os arquivos selecionados |  |
|             | Os arquivos                    |  |

Incremental

criados ou  
alterados  
desde o  
último  
backup

Diferencial

Os arquivos  
criados ou  
alterados  
desde o  
último  
backup

Diário

Os arquivos  
modificados  
ou criados  
numa data  
específica

De cópia

Todos os  
arquivos

## 11.12. Backups cíclicos – entendendo finalmente

Imagine, caro leitor, uma empresa que começou a fazer diariamente backups de seus dados. Imagine que se trata de um banco ou uma empresa de crédito imobiliário (os dados são muito importantes e não podem se perder sob hipótese alguma).

Pois é. Você acha, caro leitor, que essa empresa, no intuito de possuir cópias seguras dos dados que detém, vai gastar uma fita por dia? Em um ano de realização de backups, serão gastas 365 fitas (366 em anos bissextos, claro)! Você realmente acha que isso é inteligente?

Os ciclos evitam isso. Fazer backup de modo cíclico permite que a empresa disponha apenas de algumas fitas (em número fixo), como, digamos, sete fitas para fazer um ciclo semanal de backup. Diminuindo o número de fitas, se diminuiria, também, o custo com elas, não é?

### 11.12.1. Só é importante o último ciclo

Quando acontece um problema em um sistema de informação de uma empresa (um servidor, por exemplo), para que serve o backup? Para restaurar os dados naquele servidor de modo a que seus dados pareçam os mais atualizados possíveis em relação ao momento quando ocorreu o sinistro.

Então, se, em uma empresa, faz-se, há mais de um ano, backup em ciclos semanais, só serão importantes, para a recuperação dos dados do sistema, as fitas dos backups dessa semana (que é o último ciclo). Ao reiniciar um ciclo, tudo o que veio antes dele é totalmente desnecessário.

Os ciclos sempre reiniciam em um backup normal. Ou seja, se um ciclo de uma empresa é semanal e começa no domingo é porque o backup normal é feito justamente no domingo.

Não há como ser diferente. Os ciclos dos backups sempre começam como um backup normal. Porque ao fazer um backup normal, copia-se TUDO. E isso, imediatamente, torna o que veio antes totalmente desnecessário.

Além disso, para recuperar os dados em caso de algum problema no sistema, deve-se contar com a fita do último backup normal (o backup normal mais recente – o backup que iniciou o ciclo atual) e as demais fitas dos backups auxiliares subsequentes, se houver.

Vamos conhecer algumas das principais maneiras de se realizar backup em ciclos.

### 11.12.2. Estratégia 1: usando apenas backups normais

Esta é a forma mais simples de gerenciar os backups em uma empresa: realizar apenas backups normais, todos os dias.

Ou seja, no domingo faz-se um backup normal, na segunda-feira realiza-se outro, na terça-feira também e assim por diante. Todos os dias o sistema vai “parar” um pouco para que os dados sejam copiados da empresa em uma fita.

Lembre-se de que ao fazer um backup normal na segunda-feira, todos os arquivos selecionados serão copiados para a fita. Isso, em si, já nos permite desconsiderar a fita do domingo (seus dados não serão mais necessários depois do backup feito na segunda).



Claro que você entendeu que ao fazer um backup normal, tudo está ali. Não há necessidade de mais nenhuma fita. A fita do backup normal encerra-se em si mesma: é completa, total! As fitas que vieram antes desse último backup normal (o mais recente) são desnecessárias.

Nesse caso, qual é o ciclo? Diário! A cada novo backup normal, o ciclo é reiniciado. É “um ciclo de um backup só”, se preferir. Em resumo:

a. Só se precisa utilizar uma única fita (já que o backup de hoje desconsidera o de ontem).

b. Para recuperar o sistema em caso de sinistro, usa-se apenas a fita do último backup normal (que, se levarmos em consideração que realmente há uma única fita no sistema, é a ela mesma que estamos nos referindo).

Portanto, se o problema ocorrer em uma quinta-feira, para recuperar o sistema, basta usar a fita do backup normal que foi realizado na quarta-feira, e o sistema “voltará à vida” com os mesmos dados que tinha na quarta-feira à noite (é o mais próximo que se pode chegar dos dados atuais, não é mesmo?).

A seguir um pequeno quadro dessa estratégia de backup.

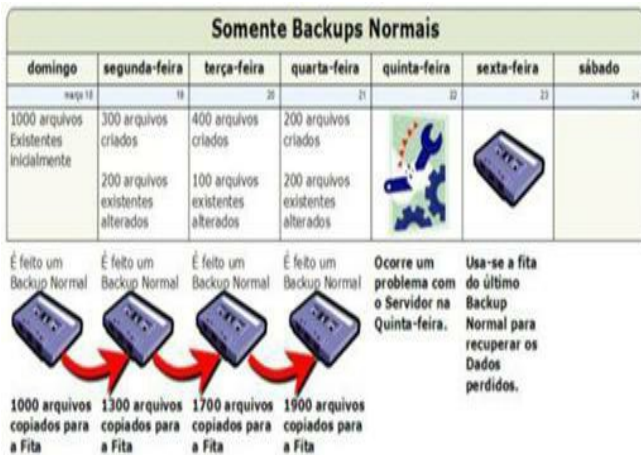


Figura 11.7 – Só backups normais.

### 11.12.3. Estratégia 2: usando backups normais + backups incrementais

Esta é muito interessante! Uma das formas mais comuns de backups exigidas em prova! Vamos imaginar que a regra usada para fazer as cópias de segurança seja:

- a. Domingo: backup normal.
- b. Demais dias da semana: backups incrementais.

Diante dessa rotina definida, temos alguns pontos a entender:

1. Quando um backup normal é realizado (no domingo, claro), tudo o que foi feito anteriormente (backups da semana anterior) é desconsiderado. Ou seja, só para lembrar, o backup normal é quem inicia o ciclo! É o backup normal que determina o reinício de um novo ciclo.
2. Com isso, podemos concluir que a fita que foi usada para o backup normal do domingo passado poderá, naturalmente, ser utilizada para o backup normal deste domingo, sem prejuízo aos dados armazenados (será o reinício de um novo ciclo).
3. O backup incremental que é realizado logo após o backup normal (ou seja, o backup incremental da segunda-feira) copiará todos os arquivos criados ou modificados desde o backup normal. Ou seja, o backup incremental da segunda-feira copiará para a fita aqueles arquivos criados ou modificados na própria segunda-feira. Depois de copiados, os arquivos são marcados (têm seu “atributo de arquivamento” desmarcado), e isso fará com que no dia seguinte nenhum arquivo “amanheça” precisando passar por um backup.
4. Os demais backups incrementais (terça, quarta etc.) copiarão os arquivos criados ou modificados desde o incremental anterior (realizado no dia anterior). Ou seja, nesse nosso exemplo, os backups incrementais só copiarão o que foi modificado ou criado naquele dia específico.
5. Os backups incrementais não acumulam os arquivos entre si, ou seja, o que for copiado para a fita na quarta-feira não incluirá, necessariamente, os arquivos que foram incluídos na fita da terça-feira. Eles não são cumulativos.
6. Em suma, cada backup incremental deve ser feito em uma fita diferente! A fita da segunda-feira não poderá ser utilizada na terça- feira; a fita da terça-feira não poderá ser utilizada na quarta-feira e assim por diante.
7. Como o backup normal reinicia tudo nos domingos, podemos desconsiderar o conteúdo das fitas da semana passada. Com isso, podemos usar as fitas dos incrementais da semana passada nos incrementais atuais.

Com isso tudo, chegamos a algumas conclusões com relação à rotina de backups que usa o backup normal + backup incremental:

- a. Será necessária uma fita para o backup normal.
- b. Será necessária uma fita para cada backup incremental (ou seja, no nosso exemplo, seis fitas para incrementais).
- c. Para recuperar o sistema em caso de algum problema, será necessário restaurar primeiramente os dados da fita do último backup normal e, em seguida, as fitas dos backups incrementais posteriores (na mesma ordem em que foram copiadas).

Portanto, tomemos como exemplo um sistema que “deu pau” na quinta-feira. Para recuperar todos os dados de modo que o sistema volte a ser como era (ou pelo menos, o mais próximo possível do que era), será necessário recuperar:

- a. O backup normal do último domingo.

- b. O backup incremental da última segunda-feira.  
 c. O backup incremental da última terça-feira.  
 d. O backup incremental da última quarta-feira (o dia anterior ao sinistro).  
 A seguir, o esquema desta rotina de backups.

| Domingos: Backup Normal ----- Outros dias: Backup Incremental                    |                                                                                   |                                                                                   |                                                                                   |                                                                                   |                                                                                                                                                    |        |
|----------------------------------------------------------------------------------|-----------------------------------------------------------------------------------|-----------------------------------------------------------------------------------|-----------------------------------------------------------------------------------|-----------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------|--------|
| domingo                                                                          | segunda-feira                                                                     | terça-feira                                                                       | quarta-feira                                                                      | quinta-feira                                                                      | sexta-feira                                                                                                                                        | sábado |
| 19                                                                               | 19                                                                                | 20                                                                                | 21                                                                                | 22                                                                                | 23                                                                                                                                                 | 24     |
| 1000 arquivos<br>Existentes<br>inicialmente                                      | 200 arquivos<br>existentes<br>alterados<br><br>300 arquivos<br>criados            | 100 arquivos<br>existentes<br>alterados<br><br>100 arquivos<br>criados            | 100 arquivos<br>existentes<br>alterados<br><br>200 arquivos<br>criados            |  |                                                                   |        |
| É feito um<br>Backup Normal                                                      | É feito um<br>Incremental                                                         | É feito um<br>Incremental                                                         | É feito um<br>Incremental                                                         | Ocorre um<br>problema com<br>o servidor na<br>Quinta-feira.                       | Usam-se, em<br>ordem, as<br>fitas do...<br>- Normal (1);<br>- Inc. (2);<br>- Inc. (3);<br>- Inc. (4);<br>para<br>recuperar os<br>dados<br>perdidos |        |
|  |  |  |  |                                                                                   |                                                                                                                                                    |        |
| 1000 arquivos<br>copiados para<br>a Fita 1                                       | 500 arquivos<br>copiados para<br>a Fita 2                                         | 200 arquivos<br>copiados para<br>a Fita 3                                         | 300 arquivos<br>copiados para<br>a Fita 4                                         |                                                                                   |                                                                                                                                                    |        |

Figura 11.8 – Backups normais + incrementais.

#### 11.12.4. Estratégia 3: usando backups normais + backups diferenciais

Olha aqui a “opção” em relação ao exemplo anterior. Trocar, simplesmente, os incrementais pelos backups diferenciais.

Vamos imaginar a mesma estratégia anterior: aos domingos, serão feitos backups normais; nos demais dias da semana, backups diferenciais. Os pontos importantes desse método são:

1. Quando o backup normal for realizado, todos os arquivos do sistema serão copiados. Sem distinção! Depois disso, todos os arquivos serão marcados como copiados (ou seja, seu “atributo de arquivamento” será desmarcado). Isso servirá para indicar quem já passou pelo backup (que, no caso, foi todo mundo).
2. Novamente, o backup normal é o início de um novo ciclo, permitindo que tudo o que se tinha antes dele seja desconsiderado.

3. Quando chegar o momento de fazer um backup na segunda-feira (um diferencial, segundo nossa estratégia), ele copiará para a fita todos os arquivos que foram criados ou alterados desde o último backup (no caso, o backup normal do domingo). Isso significa que, na fita da segunda-feira, haverá apenas arquivos que foram criados ou modificados naquela segunda-feira.

4. Como o backup é diferencial, ele não marcará ninguém (deixará o “atributo de arquivamento” inalterado, do jeito como estava). Isso faz com que esses arquivos já passem para a terça-feira desmarcados (com o “atributo” marcado), ou seja, precisando passar novamente por um backup.

5. O backup diferencial da terça-feira incluirá na fita daquele dia os arquivos criados ou modificados naquele mesmo dia (terça) e também aqueles que foram “herdados” da segunda-feira (porque não foram marcados no dia anterior). A fita da terça-feira, portanto, incluirá os arquivos da fita da segunda-feira. Depois do backup da terça, nenhum arquivo é marcado, passando para a quarta-feira precisando novamente ser backuppeados.

6. O backup diferencial da quarta-feira incluirá naquela fita os arquivos criados ou modificados na quarta, mas também incluirá tudo aquilo que veio desmarcado (com o “atributo de arquivamento” marcado) da terça-feira. Ou seja, a fita da quarta necessariamente incluirá o conteúdo da fita da terça-feira.

7. E assim por diante. Em uma sequência ininterrupta de backups diferenciais, todas as fitas são cumulativas, pois trazem, necessariamente, o conteúdo de todas as fitas dos diferenciais anteriores. Em suma, na fita da sexta-feira, estarão, sem sombra de dúvidas, os arquivos que foram copiados nas fitas da segunda-feira, terça, quarta e quinta.

Diante disso, podemos concluir que:

- a. Será necessária uma fita para os backups normais (podendo ser reutilizada a cada backup normal).
- b. Será necessária apenas uma única fita para todos os backups diferenciais (que poderá ser reutilizada no próximo diferencial).
- c. Caso aconteça algo com o sistema, para recuperar seus dados é necessário usar apenas a fita do último backup normal e, em seguida, a fita do último backup diferencial.

Vamos ver, como exemplo, um sistema que “deu bronca” na quinta-feira. (Que diazinho cheio de “mazela”, hein? As coisas só acontecem nesse dia.) Para recuperar o sistema em caso de problemas que ocasionaram perda de dados na quinta-feira, deve-se:

- a. Recuperar a fita do backup normal do último domingo (o normal mais recente, ou seja, que iniciou o ciclo atual); e
- b. Recuperar a fita do backup diferencial mais recente (o da quarta- feira, no nosso caso).

Na figura a seguir, a rotina com diferenciais.

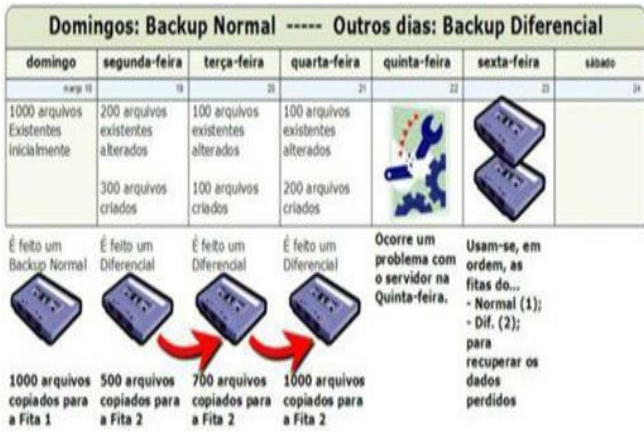


Figura 11.9 – Usando normais + diferenciais.

Apenas um aviso: para a figura anterior, que versa sobre os backups diferenciais, devemos considerar que os “arquivos existentes alterados” descritos na terça não foram copiados previamente na segunda (ou seja, são diferentes dos já copiados em um backup diferencial anterior) para poder “bater” a soma de arquivos da segunda com os alterados na terça.

O pensamento é o mesmo para os “arquivos existentes alterados” na quarta: eles são diferentes dos 700 arquivos copiados para fita na terça. Se não fosse assim, a soma não daria exatamente 1.000 arquivos no diferencial da quarta-feira.

#### 11.12.5. Estratégia 4: usando backups normais + incrementais + diários

Aqui temos uma rotina mais complexa, usada em ciclos maiores como um ciclo mensal (que é o nosso exemplo). Vamos às regras básicas da nossa estratégia:

- Dia 1<sup>o</sup> do mês: backup normal;
- Domingos: backups incrementais;
- Demais dias: backups diários.

Então, nesse caso, o ciclo se inicia sempre no dia 1<sup>o</sup> de cada mês! Nesse dia, é feito um backup normal, que copiará todos os arquivos para a fita, tornando desnecessária toda e qualquer fita que contenha backups anteriores. Vamos às conclusões:

- O backup normal inicia o ciclo, sempre!

2. Os backups diários, usados nos dias da semana, copiarão para a fita apenas os arquivos criados ou alterados naquele dia específico. Depois da cópia feita, nenhum arquivo é marcado (não há alteração no “atributo de arquivamento”).
3. Os backups incrementais dos domingos funcionarão como uma espécie de “consolidado” da semana. Como um backup incremental só copia para a fita aqueles arquivos que precisam ser copiados, ou seja, aqueles que não estão marcados (ou, em outras palavras, que possuem o “atributo de arquivamento” marcado), serão copiados para a fita do domingo todos os arquivos criados ou alterados naquela semana.
4. Se for o primeiro incremental do mês, serão copiados para a fita todos os arquivos criados ou modificados desde o backup normal do dia 1<sup>o</sup>. Não importando se foram, ou não, copiados em fitas dos diários.
5. Se não for o primeiro incremental do mês, serão copiados para a fita todos os arquivos criados ou alterados desde o backup incremental anterior (do domingo anterior).
6. Em ambos os casos (item 4 ou 5), o backup incremental servirá como “reiniciador” do ciclo semanal, pois conterá tudo o que foi alterado naquela semana (como vimos, como um “consolidado” semanal). Sendo assim, as fitas dos backups diários poderão ser reutilizadas na semana posterior (depois do incremental).

Nesse sentido, notamos a necessidade de:

- a. 1 fita para o backup normal;
- b. 5 fitas para os incrementais (uma para cada domingo do mês);
- c. 6 fitas para os diários (uma para cada dia da semana);

Entende-se, por fim, que caso haja algum problema no sistema, digamos, em uma quinta-feira, dia 17, será necessário recuperar, em ordem, os dados das fitas dos backups:

- a. Normal – realizado no dia 1<sup>o</sup>, que, no nosso caso, foi uma terça-feira.
- b. Incremental – realizado no domingo (dia 06).
- c. Incremental – realizado no domingo (dia 13).
- d. Diários – realizados nos dias 14 (segunda), 15 (terça) e 16 (quarta).

#### **11.12.6. Estratégia 5: usando backups normais + diferenciais + diários**

Essa estratégia difere muito pouco da anterior. Apenas não será necessário ter várias fitas para os backups “consolidados” dos domingos. Será necessária apenas uma única fita. Vamos às regras básicas:

- a. Dia 1<sup>o</sup> do mês: backup normal.
- b. Domingos: backups diferenciais.
- c. Demais dias: backups diários.

Vamos às primeiras conclusões:

1. O backup normal inicia o ciclo mensal (que é o nosso caso).
2. Os backups diários, usados nos dias da semana, copiarão para a fita apenas os arquivos criados ou alterados naquele dia específico. Depois da cópia feita, nenhum arquivo é marcado (não há alteração no “atributo de arquivamento”). Quanto a isso não há alteração, pois é o “jeito” de o backup diário trabalhar.

3. Os backup diferenciais dos domingos não funcionarão exatamente como o “consolidado” da semana (como o incremental) apesar de só copiarem para a fita, da mesma forma que o incremental, aqueles arquivos que precisam ser copiados, ou seja, aqueles que não estão marcados (ou, em outras palavras, que possuem o “atributo de arquivamento” marcado).

4. Justamente pelo seu funcionamento (pela diferença de funcionamento entre ele e o incremental, para ser mais preciso), o backup diferencial (não importando se é o primeiro do mês ou não) copiará sempre os arquivos criados ou modificados desde o backup normal (ou seja, desde o dia 1<sup>o</sup>).

5. Sempre o backup diferencial de um domingo incluirá, necessariamente, os arquivos contidos no backup diferencial do domingo anterior. Portanto, só há a necessidade de uma única fita para todos os diferenciais, já que ela poderá ser reutilizada.

Chegamos à conclusão, em matéria de necessidades logísticas para essa estratégia, que precisaremos de:

- 1 fita para o backup normal.
- 1 fita para os diferenciais (uma única para todos eles no mês).
- 6 fitas para os diários (uma para cada dia da semana).

Simples, não?

### 11.13. Termos e trechos “especiais” nas provas

Há alguns termos interessantes que já foram palco de polêmica em várias provas de concursos. Trago para vocês alguns deles...

1. “Tanto o backup incremental quanto o diferencial copiam os arquivos criados ou modificados desde o último backup normal ou incremental.”

Bem, sabemos que o backup incremental (e o diferencial também) são backups que só copiam quem precisa ser copiado, ou seja, só copiam os arquivos que foram criados ou alterados recentemente. Mas, exatamente “quanto” recentemente?

Os backups incrementais e diferenciais levam em consideração o estado do indicador de backup (a “marca” ou “atributo de arquivamento”), não é? Portanto, tudo o que foi alterado ou criado depois de um “backup marcador” será incluído no incremental e diferencial.

Mas o que é um “backup marcador”? É um termo que uso para descrever os backups que, depois de realizados, marcam os arquivos (ou seja, desmarcam o “atributo de arquivamento”). Esse termo você não encontrará na prova, é somente para que eu possa explicar-lhe agora.

Imagine que em um domingo foi realizado um backup marcador (seja ele um normal ou um incremental, não importa). Depois desse domingo (na manhã da segunda-feira subsequente), quantos arquivos estarão desmarcados (com o “atributo” marcado), ou seja, precisando passar por um backup? Resposta: nenhum!

Todos os arquivos presentes no sistema da empresa naquela manhã de segunda-feira estarão marcados (com o “atributo” desmarcado) porque, se não todos, aqueles que estavam precisando passar por um backup no domingo foram marcados após o backup.

Então, a expressão “tanto o backup incremental quanto o diferencial copiam os arquivos criados ou modificados desde o último backup” é um tanto incompleta, porque a palavra backup, nesse caso, pode se referir a um backup “não marcador”. A expressão certa seria, mesmo, a que

vemos no início deste tópico (1)!

2. “Backup diário copia os arquivos criados ou alterados na data corrente.”

Obviamente, sabemos que essa frase é “enganadora” porque o backup diário nos permite escolher a data exata a qual queremos fazer o backup. Mesmo retroativamente, como vimos.

Mas, pela lógica, a razão do backup diário é essa mesma! O sentido de se fazer backup diário é, realmente, copiar os arquivos alterados ou criados na data em que o backup diário está sendo realizado (daí a razão do “na data corrente”).

Essa expressão é considerada verdadeira, onde quer que apareça!

3. “Espelhamento de Disco (RAID 1) é uma forma de backup automático.”

Existe uma série de questões que tentam equivaler um ato de backup ao uso de outras técnicas (de hardware) para evitar (ou corrigir) perdas de dados. Uma delas é associar um ato de backup ao RAID 1 (conjunto de discos rígidos espelhados).

Backup é um ato deliberado de cópia de arquivos em outro disco. Backup tem de ser realizado com hora definida, é um processo realizado por alguém de carne e osso, ou por uma máquina que coloque, copie e retire as fitas.

RAID 1 é um recurso muito interessante, mas não é backup. RAID 1, como já foi visto no capítulo sobre hardware, faz com que os dados sejam copiados automaticamente para dois HDs idênticos. Se um deles falhar, o outro disco assumirá automaticamente.

E se os dois falharem? Os dados serão perdidos! Ai entra em cena o backup.

O uso de uma matriz RAID 1 dispensa a realização de backups? Nem sonhando!

Mas, como nem tudo é exatamente como a gente quer. Já houve uma dezena de questões de provas que dizem “é um backup automático, feito pela placa-mãe num disco espelhado” e traziam que RAID 1 era a resposta (paciência, né?!).

Portanto, caríssimo leitor, eu sugiro que você entenda, e defina, RAID 1 como sendo uma “espécie” de Backup... Pelo menos para as provas!

## 11.14. Palavras finais

É isso, caro leitor, viu como o assunto é pequeno? Viu como é fácil entendê-lo? Agora é partir para os testes de aptidão. Veja algumas questões que criei sobre o assunto e algumas que a Esaf e a FCC fizeram.

Continue estudando!



## 11.15. Q u e s t õ e s d e B a c k u p

### 1. O backup normal:

- I. copia todos os arquivos selecionados;
  - II. copia apenas os arquivos criados desde o último incremental;
  - III. leva em consideração a data de alteração dos arquivos para copiá-los;
  - IV. marca os arquivos, indicando que eles passaram por um backup.
- Quantas proposições são verdadeiras?

- a) 1.
- b) 2.
- c) 3.
- d) 4.
- e) 5.

### 2. Em uma rotina de backups, realizada há mais de um ano, que se baseia em backups normais e incrementais, o processo de recuperação de dados perdidos em um sinistro:

- a) será realizado usando-se apenas o último backup incremental;
- b) necessita do primeiro backup normal realizado desde o início das rotinas;
- c) precisa pelo menos do último backup normal;
- d) poderá ser realizado com sucesso sem a fita do último backup normal;
- e) necessita de todos os backups incrementais realizados desde o início da rotina de backups da empresa.

### 3. Um backup diário:

- a) não pode ser realizado em uma rotina de backup em que também é descrita a existência de backups diferenciais;
- b) copia os arquivos criados ou modificados em uma data específica (normalmente a data atual);
- c) copia os arquivos criados ou modificados desde o último backup normal;
- d) copia todos os arquivos selecionados pelo usuário;
- e) altera o “atributo de arquivamento” dos arquivos copiados por ele.

### 4. Se o “atributo de arquivamento” do arquivo Orçamento.xls foi marcado, é correto afirmar que:

- a) o arquivo Orçamento.xls foi copiado por um backup normal recentemente;
- b) o arquivo Orçamento.xls foi copiado por um backup incremental recentemente;
- c) o arquivo Orçamento.xls foi recuperado recentemente;
- d) o arquivo Orçamento.xls foi apagado recentemente;
- e) o arquivo Orçamento.xls foi alterado recentemente.

### 5. Analise a rotina de backups apresentada a seguir:

- Nos domingos são realizados backups normais;
- Nos demais dias da semana, são feitos backups diferenciais;

Em caso de problemas no computador onde os dados estão hospedados, o processo de recuperação desses dados:

- a) será realizado com o uso de apenas duas fitas: sendo uma do último backup normal e a outra do último backup diferencial posterior;
- b) será realizado com o uso de quatro fitas, caso o problema tenha ocorrido numa quinta-feira;
- c) precisará de todas as sete fitas;
- d) não poderá ser realizado sem a fita do primeiro backup normal;
- e) só poderá ser realizado depois da realização de, pelo menos, três backups normais (ou seja, depois de dois ciclos completos).

6. Deve-se tomar alguns cuidados com as informações armazenadas em um computador. Um dos cuidados mais importantes é a realização de cópias de segurança (Backup). Com relação ao backup, é correto afirmar que:

- a) o mais importante é a realização, diária, da cópia de segurança do Sistema Operacional de sua máquina;
- b) quando se realiza uma cópia de segurança do conteúdo de uma pasta que se encontra no disco principal de uma máquina, por exemplo, disco C:, para uma pasta denominada BACKUP, no mesmo disco, a recuperação total dos dados dessa pasta BACKUP é possível utilizando-se apenas o Windows e suas ferramentas básicas, mesmo se o referido disco for formatado;
- c) deve ser feita uma cópia de segurança dos arquivos temporários do Windows sempre que se enviar um e-mail;
- d) um backup incremental é aquele que copia somente os arquivos criados ou alterados desde o último backup normal ou incremental;
- e) uma cópia só pode ser considerada segura se for realizada em um disquete.

7. Os tipos de backups determinam quais dados sofrem a cópia de segurança e a forma como ela deve ser feita. Com relação a este assunto é correto afirmar que:

- a) o backup incremental deve ser feito sempre antes de um backup normal;
- b) o backup normal deve ser feito sempre após um backup diferencial e só deve ser descartado após o próximo backup incremental;
- c) o uso de um backup normal diário dispensa o uso de um backup incremental semanal;
- d) o uso de um backup diferencial após um backup normal pode danificar todo o sistema de backup de uma empresa se, após a sua realização, não for feito um backup incremental;
- e) a principal diferença entre os backups normal, incremental e diferencial está no sistema de fitas utilizado para armazená-los.

8. Uma forma de proteger os dados de uma organização contra perdas acidentais é a realização periódica do backup desses dados de uma forma bem planejada. Entre os tipos de backup, no incremental:

- a) é feito o backup dos arquivos selecionados ou indicados pelo usuário somente se eles não tiverem sido marcados como copiados (participado do último backup) ou se tiverem sido

alterados, marcando-os como copiados (marca que indica que participaram do último backup);

b) é feito o backup de todos os arquivos selecionados ou indicados pelo usuário, independentemente de estarem marcados como copiados (participado do último backup), marcando-os como copiados (marca que indica que participaram do último backup);

c) é feito o backup de todos os arquivos selecionados ou indicados pelo usuário, independentemente de estarem marcados como copiados, mas nenhum é marcado como copiado (marca que indica que participaram do último backup);

d) é feito o backup dos arquivos selecionados ou indicados pelo usuário somente se eles não tiverem sido marcados como copiados (participado do último backup) ou se tiverem sido alterados, mas nenhum é marcado como copiado (marca que indica que participaram do último backup);

e) é feito o backup apenas dos arquivos selecionados ou indicados pelo usuário que tiverem sido alterados na data corrente, mas não marca nenhum como copiado (marca que indica que participaram do último backup).

9. Um backup \_\_\_\_\_ (I) copia apenas os arquivos criados ou alterados \_\_\_\_\_ (II). Completam as lacunas:

a) (I) – Incremental; (II) – desde o último Backup Diferencial ou Incremental;

b) (I) – Diário; (II) – desde o último Backup Normal ou Incremental;

c) (I) – Incremental; (II) – desde o último Backup Diferencial ou Normal;

d) (I) – de Cópia; (II) – desde o último problema ocorrido;

e) (I) – Diferencial; (II) – desde o último Backup Normal ou Incremental.

#### 12.1. Sistemas numéricos

Chegamos a um assunto que não aparece realmente em todos os concursos: a aritmética computacional (genericamente conhecida como “operações com bases numéricas usadas em computação”). Isso quer dizer, simplesmente, trabalhar com matemática diferente daquela com a qual estamos acostumados.

A mais comum operação a ser exigida neste assunto é a conversão de números entre bases diferentes. Para que possamos ser capazes de converter uma base numérica para outra, devemos conhecer o sentido de cada uma delas. É o que vamos fazer agora.

Um sistema numérico é um conjunto de regras de escrita de números baseado em uma quantidade definida de caracteres. O sistema numérico que usamos é conhecido como sistema decimal, pois utiliza 10 símbolos diferentes (dígitos, caracteres, como queira chamar). A base de um sistema numérico é exatamente o número de caracteres que o forma. Em informática, nos deparamos, além do sistema decimal, com os sistemas numéricos binário, octal e hexadecimal.

Segue uma pequena explicação de cada um deles, incluindo as listagens de seus números componentes. Mais adiante trataremos de cada um com a devida atenção, apresentando, inclusive, as principais operações aritméticas usadas em cada um dos sistemas, bem como as técnicas de conversão de um sistema para outro.

- **Sistema Decimal:** usa dez dígitos (0, 1, 2, 3, 4, 5, 6, 7, 8, 9). É o sistema numérico usado por nossa matemática. Desde pequenos, nos acostumamos a fazer cálculos com esses números, sempre baseando-nos nesse sistema. Esse sistema tem base 10.
- **Sistema Binário:** por natureza, é o sistema usado pelos computadores. Utiliza apenas dois dígitos (0 e 1). Por seu formato, o sistema binário admite certas operações que não conseguimos usar em nosso sistema decimal, além, é claro, de não aceitar, também, certas operações que costumamos realizar. Esse sistema tem base 2.
- **Sistema Octal:** não entendo por que inventaram esse! É um sistema numérico baseado em oito dígitos (0, 1, 2, 3, 4, 5, 6, 7). Acho que foi só para complicar a nossa vida mesmo! Adivinha? Esse sistema tem base 8.
- **Sistema Hexadecimal:** se o outro era estranho, imagina esse! O sistema hexadecimal é formado por 16 dígitos (0, 1, 2, 3, 4, 5, 6, 7, 8, 9, A, B, C, D, E, F). Normalmente usado para representar, de forma mais fácil, os números binários grandes (mais adiante mostrarei as conversões de um sistema para o outro). Esse sistema tem base 16 (claro!).

“João, mas qual é a utilidade desse ‘trem’?”

Caro leitor, esse “trem” é usado por muitos assuntos dentro da informática; afinal, as informações que os nossos computadores conseguem manipular são representadas por esse tipo (formato) de dado.

Um exemplo simples é o endereço IP e as operações com máscaras de sub-rede. (Se você chegou aqui por causa do assunto de redes, naquele ponto exato em que te mandei vir para cá, ótimo... Continue lendo o capítulo!)

## 12.2. Como é formado um número

Vamos fazer uma análise na maneira com que um número qualquer é formado em nosso sistema decimal. Começando, é claro, bem do início. Todos os números (inteiros positivos) que conhecemos como números naturais fazem parte de uma sequência ininterrupta e infinita (isso, você sabe desde o primário).

0, 1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 15, 16...

Note que, quando acabam os caracteres disponíveis para continuar a sequência (acabou exatamente no 9), é necessário inserir um novo “espaço” para um algarismo. Por convenção, esse espaço foi preenchido à esquerda (antes) do primeiro algarismo. Na verdade, esse “espaço” já era usado por um dígito “0” (zero), que, quando está à esquerda, não é mostrado. (Até aqui, alguma novidade? Acho que não.)

Analisemos o número 9342 como exemplo. No primário, a “tia” provavelmente diria “esse número é formado por 9 milhares, 3 centenas, 4 dezenas e 2 unidades”, não é mesmo? Pois é exatamente dessa forma que devemos analisar esse número!

Um número é formado pela multiplicação de cada dígito pela base do sistema elevada a uma potência crescente em cada “espaço” de dígito. Esse crescimento acontece da direita para a esquerda, ou seja, começando na casa das unidades.

“O quê? Não entendi nada, João!”

Nem eu, caro leitor. Vou tentar explicar melhor.

A base do nosso sistema numérico é 10. Então, todos os números formados em nosso sistema são baseados em potências de 10. Veja o exemplo a seguir para o mesmo número 9342:

|                               |                 |                |  |
|-------------------------------|-----------------|----------------|--|
| O dígito...                   | 9               | 4              |  |
| O valor da casa...            | $10^3$          | $10^2$         |  |
| Que significa...              | 1000            | 100            |  |
| O dígito x o valor da casa... | $9 \times 1000$ | $4 \times 100$ |  |

Que  
resulta  
em...

9000

400

Por fim,  
somando,  
dá...

$900 + 400 + 30 + 2$

A “casa” das unidades (corresponde ao primeiro dígito, da direita para a esquerda) vale  $10^0$ , ou seja, vale 1. Ela é preenchida por um número que será multiplicado por seu valor. A “casa” das dezenas, por sua vez, vale  $10^1$ , ou seja, 10. Qualquer valor presente nesta casa será multiplicado por 10. A terceira casa vale  $10^2$ , ou seja, 100. Qualquer valor nessa casa será multiplicado por 100 e assim por diante.

Em todos os sistemas numéricos, os números evoluem (ou seja, são incrementados) um em um na primeira casa (unidades). Quando atingimos o último valor (dígito) possível nessa casa, reiniciamos o ciclo da primeira casa, mas incrementando a segunda casa. É assim que os números 0, 1, 2, 3, 4, 5, 6, 7, 8 e 9 viram 10 (9 vira 0, mas adiciona 1 à casa da esquerda). É assim, também, que o 99 (último número de dois dígitos) vira 100 (reinicia-se o ciclo nas duas primeiras casas, e adiciona-se 1 à terceira).

**Binário**

**Octal**

**Decimal**

0

0

0

1

1

1

10

2

2

|       |    |    |  |
|-------|----|----|--|
| 11    | 3  | 3  |  |
| 100   | 4  | 4  |  |
| 101   | 5  | 5  |  |
| 110   | 6  | 6  |  |
| 111   | 7  | 7  |  |
| 1000  | 10 | 8  |  |
| 1001  | 11 | 9  |  |
| 1010  | 12 | 10 |  |
| 1011  | 13 | 11 |  |
| 1100  | 14 | 12 |  |
| 1101  | 15 | 13 |  |
| 1110  | 16 | 14 |  |
| 1111  | 17 | 15 |  |
| 10000 | 20 | 16 |  |
| 10001 | 21 | 17 |  |
|       |    |    |  |

|       |     |     |
|-------|-----|-----|
| 10010 | 22  | 18  |
| 10011 | 23  | 19  |
| 10100 | 24  | 20  |
| 10101 | 25  | 21  |
| 10110 | 26  | 22  |
| ...   | ... | ... |

### 12.2.1. Entendendo alguns detalhes importantes

Vamos analisar alguns detalhes interessantes nos números. Detalhes esses que acontecem em todas as bases numéricas; portanto, auxiliarão você no processo de conversão.

1. Sempre o primeiro número construído com um certo número de algarismos é uma potência da base numérica em si! E mais! É uma potência da base numérica elevada ao número de algarismos subtraído de 1.

Vamos ver isso em um exemplo da base decimal: qual é o primeiro número decimal escrito com dois algarismos? A resposta é 10 (dez) – pois bem, 10 é simplesmente  $10^1$ . Note que há dois algarismos e que o expoente é 2-1 (1).

De novo: qual é o primeiro número, no conjunto dos números naturais, escrito com cinco algarismos? A resposta é 10.000 (dez mil). De novo, 10.000 é  $10^4$  – note novamente que o expoente é 4 (5-1).

Só mais uma vez: qual é o primeiro número, dentre os números positivos, a ser escrito com nove dígitos? A resposta é 100.000.000 (cem milhões). Pois bem, cem milhões é o equivalente a  $10^8$  – sim, mais uma vez, o expoente é o número de dígitos menos um ( $8 = 9 - 1$ ).

Então, lembre-se: o primeiro número com X algarismos é sempre  $10^{X-1}$  (se considerarmos o cálculo na base decimal – base 10).

Atenção: essa ideia só não vale para o  $10^0$ , que não é equivalente ao primeiro número escrito com um algarismo (0), mas 1.

2. O primeiro número escrito com uma certa quantidade de algarismos (digamos X algarismos) é sempre precedido do último número possível escrito com X-1 algarismos.

Essa dói de tão óbvia! O primeiro número escrito com quatro algarismos é 1.000 ( $10^3$ ). O número que o precede é 999 (justamente o último número possível de escrever com três



algarismos).

Novamente, o primeiro número natural escrito com sete algarismos é 1.000.000 (um milhão –  $10^6$ ). O número que o precede imediatamente é 999.999 (que é o último número que se pode escrever com seis algarismos).

Então podemos dizer assim: o último número com  $X$  algarismos é sempre  $10^X - 1$ . (Se considerarmos a base decimal! Se fosse a binária, seria  $2^X - 1$ .)

3 O primeiro número que se pode escrever com  $X$  algarismos é sempre formado de um primeiro algarismo 1 seguido de  $X - 1$  algarismos 0 (em todas as bases).

Portanto, o primeiro número que se pode escrever usando cinco algarismos é 10.000 (dez mil). O primeiro número que se pode escrever com três algarismos é 100 (cem). E, finalmente, o primeiro número possível de escrever-se com oito algarismos é 10.000.000 (dez milhões).

“Ei João, isso é realmente necessário?”

Caro leitor, sei que está bastante óbvio – eu até diria que é um “assunto idiota”. Mas você verá que isso facilitará sua vida mais adiante. Eu prometo!

Vamos agora aos processos oficiais de conversão de bases. Depois mostrarei que há métodos menos convencionais que seguem a lógica aqui apresentada.

## 12.3. Processo de conversão de bases

### 12.3.1. Da base decimal para qualquer outra base

Para converter um número escrito no sistema decimal (base 10) para qualquer outro sistema, utilizamos um processo de sucessivas divisões tendo, como divisor, o número da base para a qual se vai converter o número.

Que tal darmos um exemplo convertendo da base 10, para, digamos... a base 10!

“Mas, João... Converter da base 10 para a base 10 é como traduzir de português para português! Não tem sentido, não chega a ser uma tradução!”

Eu sei, caro leitor! Mas é só para ilustrar o processo como um todo. Acompanhe-me na conversão do número 2584 da base decimal (base 10) para a base decimal (base 10)! Pode parecer estranha e até mesmo tola a minha sugestão, mas vai ajudar, com certeza, para que você consiga entender todo o processo.

Então, para converter o número 2584 para a base 10, devemos dividi-lo várias vezes pelo número correspondente à base numérica de destino. Ou seja, se o queremos convertido para a base 10, o divisor dessas sucessivas divisões será sempre 10. Vamos à primeira divisão:

1ª Divisão

2584  
(4)

|     |
|-----|
| 10  |
| 258 |

Começamos dividindo 2584 por 10. O resultado da divisão é 258, com resto 4. (É uma mania antiga minha colocar o resto entre parênteses. Por favor, perdoe-me se você não faz assim!

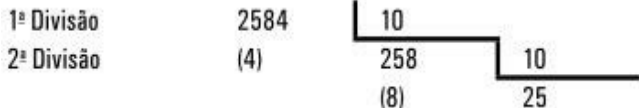
Afinal, “cada doido com sua mania”.)

“João, será necessário mesmo o resto?”

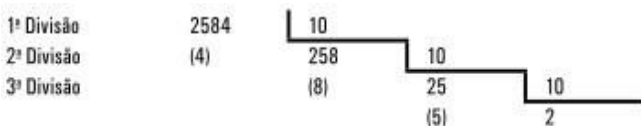
E como, caro leitor! O resto será “a alma do negócio”.

“E é só isso? Já terminamos?”

Não. As sucessivas divisões só terminam quando o quociente for menor que o divisor! Ou seja, quando o resultado for menor que 10. Então, como o resultado foi 258, ainda temos de continuar...



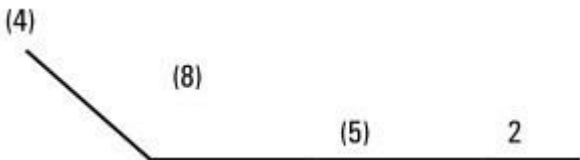
Note que foi necessário fazer uma segunda divisão, que, por sinal, ainda não foi suficiente. É necessário continuar... Vamos continuar **até o quociente ser menor que o divisor!**



Foram necessárias três divisões sucessivas para que o resultado (quociente) fosse menor que o divisor (10). Com isso nossa “conversão” chegou ao fim. Agora leia o número convertido, começando no último quociente, seguido de todos os restos sucessivamente do último até o primeiro.

“Como é que é? Não entendi nada!”

Ahhh... Desculpe, leitor. Se você desenhar a seta que lhe mostrarei agora, ficará mais fácil de entender!



Notou? 2584, que é o nosso número de origem, na base 10, quando convertido para a mesma base (decimal) vale (pasmê!) 2584! É só ler o último quociente, seguido de todos os restos, do último para o primeiro. Fácil, não? Vamos lá, não resista à tentação: teste isso com vários outros números que você quiser.

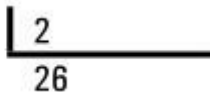
Vamos agora fazer um teste de conversão real! Vamos converter um número decimal qualquer para a base binária (base 2). O segredo da mudança é usar o valor da base de destino (2) como divisor, em vez do 10 que usamos no exemplo anterior.

Vamos dar como exemplo o número 53 na base 10. Vamos convertê-lo para a base binária:

1ª Divisão

53

(1)



Começamos dividindo o 53 por 2, obtemos 26 no quociente e 1 no resto. Não paramos por aqui, como sabemos.

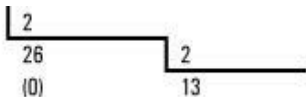
Nossa segunda divisão, por exemplo, é do 26 por 2. Vai dar resto 0! Observe:

1ª Divisão

53

2ª Divisão

(1)



Ainda temos de continuar até que o quociente seja menor que 2 (que é o divisor). Vamos agora até o final...

|            |     |     |     |     |     |   |  |
|------------|-----|-----|-----|-----|-----|---|--|
| 1ª Divisão | 53  | 2   |     |     |     |   |  |
| 2ª Divisão | (1) | 26  | 2   |     |     |   |  |
| 3ª Divisão |     | (0) | 13  | 2   |     |   |  |
| 4ª Divisão |     |     | (1) | 6   | 2   |   |  |
| 5ª Divisão |     |     |     | (0) | 3   | 2 |  |
|            |     |     |     |     | (1) | 1 |  |

Foram necessárias cinco divisões sucessivas por 2 para saber que o número **53 da base 10** é escrito como **110101 na base 2**.

Outro exemplo, para não ficar muito difícil entender: vamos converter o número 18 da base 10 para a base 2:

|            |     |     |     |     |   |  |
|------------|-----|-----|-----|-----|---|--|
| 1ª Divisão | 18  | 2   |     |     |   |  |
| 2ª Divisão | (0) | 9   | 2   |     |   |  |
| 3ª Divisão |     | (1) | 4   | 2   |   |  |
| 4ª Divisão |     |     | (0) | 2   | 2 |  |
|            |     |     |     | (0) | 1 |  |

O número **decimal 18** é representado como **10010 na base 2**.

Vamos apenas pegar um número maior para testarmos ainda mais o nosso novo conhecimento: que tal 345 em decimal sendo convertido para binário?



|            |     |     |     |   |
|------------|-----|-----|-----|---|
| 1ª Divisão | 738 | 8   |     |   |
| 2ª Divisão | (2) | 92  | 8   |   |
| 3ª Divisão |     | (4) | 11  | 8 |
|            |     |     | (3) | 1 |

Segundo nosso processo de sucessivas divisões, o número **738 da base 10** é equivalente, **na base 8, ao número 1342**.

Agora, um último cálculo de conversão para a base 8. Desta vez, vamos fazer a conversão de um número maior, que precise de mais divisões sucessivas. Convertamos o número 11273 decimal para a base octal.

|            |       |      |     |     |   |
|------------|-------|------|-----|-----|---|
| 1ª Divisão | 11273 | 8    |     |     |   |
| 2ª Divisão | (1)   | 1409 | 8   |     |   |
| 3ª Divisão |       | (1)  | 176 | 8   |   |
| 4ª Divisão |       |      | (0) | 22  | 8 |
|            |       |      |     | (6) | 2 |

Então o número **11273 decimal** é equivalente ao número **26011 na base octal**.

Estudadas as conversões para as bases 2 e 8, analisemos, agora, o processo de conversão para a base hexadecimal (base 16)! A ideia é a mesma, visto que todas as conversões que têm como origem a base 10 são feitas pelo recurso das sucessivas divisões pela base de destino. A única coisa diferente é que o divisor será 16!

Basta apenas lembrar-se de um detalhe, caro leitor: quando o resto da divisão for maior que 9 (nove), não usaremos os números decimais que conhecemos (10, 11, 12, 13, 14 ou 15), mas sim, os caracteres usados na base hexadecimal para representá-los: A, B, C, D, E ou F (respectivamente).

Ou seja, se um dos restos for 12, usaremos a letra C, em vez do valor 12, porque o 12 decimal é escrito como C em hexadecimal.

Vamos prosseguir com a conversão de um número maior, como 23354, da base 10 para a base 16 (sistema hexadecimal).

|            |        |      |        |    |
|------------|--------|------|--------|----|
| 1ª Divisão | 23354  | 16   |        |    |
| 2ª Divisão | 10 (A) | 1459 | 16     |    |
| 3ª Divisão |        | 3    | 91     | 16 |
|            |        |      | 11 (B) | 5  |

Como resultado, o número **23354 da base 10** é equivalente ao número **5B3A em hexadecimal (base 16)**.

Vamos ver mais um exemplo para podermos prosseguir com o assunto, OK? Vamos converter o número 64202 (base 10) para a base 16.

|            |        |        |        |        |
|------------|--------|--------|--------|--------|
| 1ª Divisão | 64202  | 16     |        |        |
| 2ª Divisão | 10 (A) | 4012   | 16     |        |
| 3ª Divisão |        | 12 (C) | 250    | 16     |
|            |        |        | 10 (A) | 15 (F) |

Pronto! O número **64202** do sistema numérico decimal é escrito, no sistema hexadecimal como **FACA**. (Tudo bem! Foi uma piada infame, mas é melhor que converter o número **12.237.514**.) Tente fazer diversas outras conversões! É muito fácil, como você pôde perceber!

#### 12.3.1.1. Exemplo prático da conversão da base decimal (endereço IP)

Quer entender por que você faria isso? Quer entender em que casos você usaria a conversão de números decimais para qualquer outra base? Vou mostrar.

O principal exemplo que me vem à mente é a conversão de um endereço IP, normalmente dado em decimais (separados por pontos), como visto no capítulo sobre redes. Sabemos que um endereço IP é binário (aliás, como tudo na informática), mas escrevemos o endereço IP na forma de quatro números decimais separados por pontos.

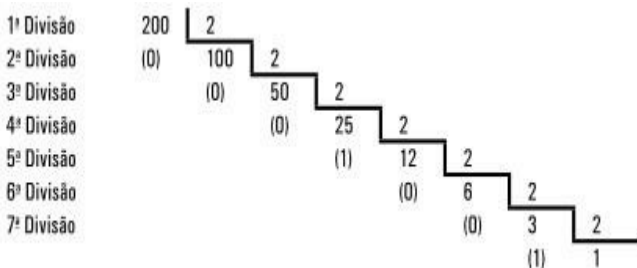
Então, o endereço 200.231.89.54 é só uma “ilusão”, uma “forma mais fácil” de representar um verdadeiro endereço formado por 32 bits (32 dígitos binários, ou seja, 32 zeros e uns). Note, caro leitor, que cada número decimal desse endereço será representado com 8 bits, ou seja, cada número desses, chamado de octeto, será, quando convertido, representado, necessariamente, por oito dígitos binários.

(Se você não está entendendo este tópico, leitor, sugiro que leia o capítulo de redes de computadores antes de prosseguir. Porém, se o assunto de redes não está no seu edital, então não

se preocupe em entender endereço IP, apenas deleite-se com as conversões!)

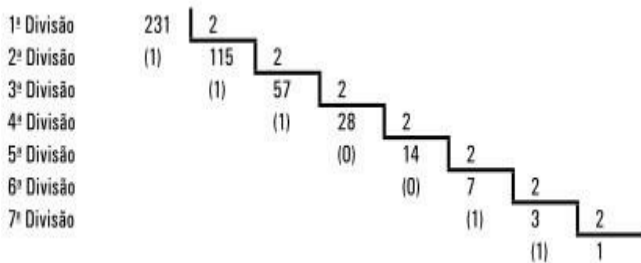
Vamos fazer a conversão?

Vamos começar com o primeiro octeto do endereço IP do exemplo: 200!



Chegamos aonde queríamos: o número 200, apresentado no primeiro octeto do endereço IP do exemplo, vale, em binário, 11001000. Note que este número já tem oito bits!

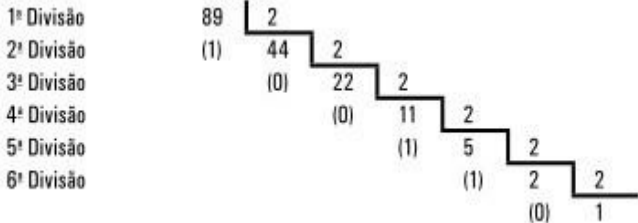
Vamos ao próximo octeto: 231.



O número 231, que é o valor que aparece no segundo octeto do endereço IP que estamos analisando, vale 11100111 em binário (também 8 bits).

Vamos agora ao terceiro octeto: 89.

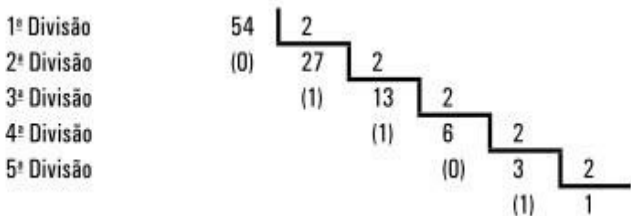




Preste atenção agora! O número 89 é representado em binário como 1011001 (tem apenas sete bits). Isso seria o suficiente se estivéssemos convertendo-o sem compromisso, apenas para saber quanto ele vale em binário. Mas no caso de um endereço IP, cada octeto do endereço tem de ser necessariamente representado com oito bits!

No caso, como o número 89 tem menos bits que isso, adicionamos vários 0 (zero) à esquerda dos bits atuais do resultado convertido, até que tenhamos oito bits! No caso do número 89, basta adicionar um único 0 (zero) à esquerda, fazendo o octeto em questão ser escrito como 01011001.

Vamos fazer a última conversão: o octeto final... 54.



O número 54 em binário é 110110. Para ser representado realmente como um octeto (octeto lembra oito, né?), ele tem de ser representado com oito bits, logo, 00110110.

Para finalizar nosso exemplo prático de conversão, temos que o endereço IP 200.231.89.54 é representado em binário como 11001000.11100111.01011001.00110110.

E como já vimos que esses pontos não existem, o endereço IP em questão é, na verdade, assim: 11001000111001110101100100110110.

### 12.3.2. De qualquer outra base para a base decimal

Para converter de qualquer outra base para a base decimal, usaremos a técnica que divide um número a ser convertido em “casas”. Atribuindo a cada casa, da direita para a esquerda, um valor que é potência da base da qual se está convertendo. O valor do expoente dessa potência é incrementado em um a cada casa...

“Hein? Agora lascou tudo!”

Calma, leitor. Esse processo já foi visto no tópico “Como é formado um número”, no início deste capítulo. O processo de conversão se dá com a multiplicação de cada algarismo do número inicial pelo valor da casa. Vamos fazer um processo passo a passo.

Começemos com o objetivo: queremos converter o número 1011010 (que é binário) para a base 10. É necessário saber o que será convertido por dois motivos: saber quantos dígitos temos (para saber quantas casas serão necessárias – no caso são sete) e saber de que base vamos converter (base de origem – no caso, base 2).

Vamos começar preparando a “tabela” que receberá os dígitos (sete “casas”):

|                |   |   |   |  |
|----------------|---|---|---|--|
| O<br>dígito... | 1 | 0 | 1 |  |
|----------------|---|---|---|--|

Esse título “O dígito...” não é necessário! Eu o coloquei para que, na hora de lermos a tabela, tenhamos certeza do que cada linha significa. O importante neste primeiro passo é colocar sete espaços para dígitos (sete “casinhas”) – um para cada dígito do número de origem.

O segundo passo é preencher o valor de cada casinha. Eis um procedimento muito fácil: preencha todos os espaços do valor das casinhas com o número da base de origem. Todas as casas!

“Você está dizendo que é para preencher tudo com 2, já que a base do número de origem é 2?”

Sim, leitor! Precisamente! Vai ficar assim!

|                    |   |   |   |   |
|--------------------|---|---|---|---|
| O<br>dígito...     | 1 | 0 | 1 |   |
| O valor<br>da casa | 2 | 2 | 2 | 2 |

é...

“Ei! Mas assim todas as casas terão o mesmo valor!”

É porque não terminou! Você vai, agora, colocar um expoente em cada um desses 2. O expoente é colocado da direita para a esquerda (da casa das unidades para a mais significativa) começando com 0 (zero) e incrementando em um a cada casinha. Vai ficar assim!

|                    |       |       |       |  |
|--------------------|-------|-------|-------|--|
| O dígito...        | 1     | 0     | 1     |  |
| O valor da casa... | $2^6$ | $2^5$ | $2^4$ |  |

“Beleza, João... E agora?”

Agora, certifique-se de saber quanto vale cada potência dessas, ou seja, saiba o que realmente significa o valor de cada casinha. (resolva as potências). Eu faria assim:

|                    |       |       |       |  |
|--------------------|-------|-------|-------|--|
| O dígito...        | 1     | 0     | 1     |  |
| O valor da casa... | $2^6$ | $2^5$ | $2^4$ |  |
| Que significa      | 64    | 32    | 16    |  |

“Preciso realmente manter essas três linhas aí? Não poderia simplesmente já preencher com os valores reais – 1, 2, 4, 8...?”

Não, você não precisa fazer as três linhas! Se já souber quais são os valores em questão, pode fazer direto, como eu, aliás, faço. Mas é necessário certo costume com esse processo! Um dia, tenho certeza, você fará cálculos de conversão de uma maneira bem reduzida! Veja uma tabela sem as linhas das potências:

|                          |    |    |    |  |
|--------------------------|----|----|----|--|
| O<br>dígito...           | 1  | 0  | 1  |  |
| O valor<br>da<br>casa... | 64 | 32 | 16 |  |

Mas vamos continuar com o processo da construção da tabela. Agora, você precisa saber quais valores devem ser somados para que se obtenha o valor real em decimal. Vamos multiplicar o valor do dígito pelo valor de sua casa.

|                             |       |       |  |
|-----------------------------|-------|-------|--|
| O dígito...                 | 1     | 0     |  |
| É<br>multiplicado<br>por... | $2^6$ | $2^5$ |  |
| Que<br>significa...         | 64    | 32    |  |

Ou seja...

1

0

x

x

64

32

Que resulta  
em...

64

0

Pronto! Chegamos ao valor que cada casa tem, levando em consideração o valor do dígito que nela “habita”. Claro que você já deve ter notado que só precisa levar em consideração as casas onde há 1, pois no binário, ou há 1, ou há 0 (e 0 – zero – vai resultar no produto 0 no valor da casa).

Pronto, para descobrir o número decimal correspondente ao número binário 1011010, basta somar os valores obtidos após a multiplicação:

$$64 + 16 + 8 + 2 = 90$$

Logo, podemos concluir que o número binário 1011010 significa 90 na base decimal. Um jeito simples de fazer a conversão que, inclusive, serve para todas as demais bases numéricas.

Mas, na base binária, podemos utilizar algumas dicas:

- Se o último dígito for 0 (zero), o número convertido para decimal é par. E, obviamente, se o último algarismo for 1 (um), o número convertido é ímpar (mas isso você já notou, não é?). Isso é só para você “tirar a prova dos nove” depois de ter convertido.
- Para fazer uma conversão rápida, faça isso: 1101 é  $8+4+1$ , que resulta em 13. (Veja que só considerei as “casas” onde há números 1, ou seja, só somei  $2^3$  com  $2^2$  com  $2^0$ . Nem precisei fazer a tabela toda.)
- De novo: 10010 significa  $16+2$  (ou seja, 18), pois só existem números 1 nas casas  $2^4$  e  $2^1$ .
- Somente mais uma: 101100 vale 44 ( $32+8+4$ ), pois só há números 1 nas casas referentes às bases  $2^5$ ,  $2^3$  e  $2^2$ .
- Esse último parágrafo estava reservado para outra dica, mas eu a esqueci...

Vamos fazer só mais uma conversão de binário para decimal: que tal converter o 11010010 para decimal? Vamos fazer do jeito mais rápido...

1

1

0

1

0

128

64

32

16

8

Somando os valores das casas onde há dígitos 1:

$$128 + 64 + 16 + 2 = 210$$

Portanto, o número 11010010 em binário é equivalente ao número 210 em decimal.

Vale salientar, porém, que esse método rápido só serve em binário (porque só há dígitos 1 e 0) – como os 0 resultarão em 0 na multiplicação, eles são desconsiderados. Como os 1 são elemento neutro da multiplicação (resultam sempre no valor com que são multiplicados), dá para usar apenas o valor da casa diretamente.

Vamos agora para uma conversão diferente: um número octal (base 8) para a base decimal (base 10). Convertamos o número 3452 da base 8 para a base 10:

|                               |                |               |  |
|-------------------------------|----------------|---------------|--|
| O dígito...                   | 3              | 4             |  |
| O valor da casa...            | $8^3$          | $8^2$         |  |
| Que significa...              | 512            | 64            |  |
| O dígito x o valor da casa... | $3 \times 512$ | $4 \times 64$ |  |
| Que resulta em...             | 1536           | 256           |  |
| Por fim,                      |                |               |  |

somando,  
dá...

$$1536 + 256 + 40 +$$

Com isso, atingimos 1834 na base decimal equivalente ao número 3452 na base 8.

Vamos agora pegar um número um pouco maior. Só para exercitar um pouco mais, vamos analisar o número octal 76510:

O dígito...

7

6

O valor  
da casa...

$8^4$

$8^3$

Que  
significa...

4096

512

Dígito x  
casa

7 x  
4096

6 x  
512

Que  
resulta  
em...

28672

3072

Por fim,  
somando,

$$28672 + 3072 + 32$$

dá...

Em resumo: o número octal 76510 é equivalente a 32072 no sistema decimal.

Agora vamos começar a realizar conversões da base hexadecimal para a base decimal. Tomemos como exemplo um número a princípio bem simples, o número 2A1C:

|                               |                    |                    |
|-------------------------------|--------------------|--------------------|
| O dígito...                   | 2                  | A <sub>(10)</sub>  |
| É<br>multiplicado<br>por...   | $16^3$             | $16^2$             |
| Que<br>significa...           | 4096               | 256                |
| Ou seja...                    | $2 \times$<br>4096 | $10 \times$<br>256 |
| Que resulta<br>em...          | 8192               | 2560               |
| Por fim,<br>somando,<br>dá... | $8192 + 2560 + 1$  |                    |



Portanto, o número  $2A1C_{(16)}$  (outra forma de se referir a um número hexadecimal) é a mesma coisa que o número  $10780_{(10)}$ . (Esse índice entre parênteses informa a base do número, como você já deve ter percebido!)

Vamos testar, só para finalizar, converter o número  $7DE_{(16)}$  para a base 10:

|                               |                   |                   |
|-------------------------------|-------------------|-------------------|
| O dígito...                   | 7                 | D<br>(13)         |
| É<br>multiplicado<br>por...   | $16^2$            | $16^1$            |
| Que<br>significa...           | 256               | 16                |
| Ou seja...                    | $7 \times$<br>256 | $13 \times$<br>16 |
| Que resulta<br>em...          | 1792              | 208               |
| Por fim,<br>somando,<br>dá... | $1792 + 208 + 14$ |                   |

Ufa! Enfim, o número  $7DE_{(16)}$  é a mesma coisa que  $2014_{(10)}$ .

Espero que você tenha entendido como realizar tais conversões. Teste seus conhecimentos nas questões que acompanham este capítulo para ver como se sai!

### 12.3.3. Da base binária para a octal (e vice-versa)

Esta conversão não tem nenhuma relação com a nossa matemática, visto que utilizamos a base decimal. Mas o processo, em si, de conversão octal-binário e binário-octal é muito simples, requer apenas que comparemos os dígitos octais com as combinações binárias, segundo esta tabela:

| Binário | Octal |
|---------|-------|
| 000     | 0     |
| 001     | 1     |
| 010     | 2     |
| 011     | 3     |
| 100     | 4     |
| 101     | 5     |
| 110     | 6     |
| 111     | 7     |

Note que cada dígito do sistema octal equivale a três dígitos do sistema binário. Portanto, para responder a uma questão que exige uma conversão binário-octal, separe os números binários em grupos de três dígitos e cada grupo desses corresponderá a um dígito octal. Tomemos como exemplo o número 101111011001 da base 2:

|                             |
|-----------------------------|
| 101111011001 <sub>(2)</sub> |
|-----------------------------|

|     |     |     |     |
|-----|-----|-----|-----|
| 101 | 111 | 011 | 001 |
|-----|-----|-----|-----|

|   |   |   |   |
|---|---|---|---|
| 5 | 7 | 3 | 1 |
|---|---|---|---|

|                                |
|--------------------------------|
| Resultado: 5731 <sub>(8)</sub> |
|--------------------------------|

Portanto, o número  $101111011001_{(2)}$  é representado como  $5731_{(8)}$ .

“Vou ter de decorar essa tabela com as oito combinações?”

Não é necessário, caro leitor, mas com certeza é mais rápido! Se você não deseja decorar, terá simplesmente de usar a técnica da conversão de binário para decimal em cada grupo de três dígitos binários (como  $101_{(2)}$ , que é, na verdade,  $2^2+2^0$ , ou seja,  $4+1$ , que resulta em  $5$  – logo,  $101_{(2)}$  é  $5_{(10)}$ , que também é  $5_{(8)}$ ).

“João, esse número binário que você deu no exemplo tem exatamente 12 dígitos, por isso deu para dividir em quatro grupos de 3 bits. E se o número não apresentar exatamente uma quantidade de dígitos que possa ser dividida por 3?”

Simples: adicione os zeros necessários no início do número e faça a divisão. Claro, para isso você terá de começar a preencher a tabela da direita para a esquerda (ou seja, da “casa” das unidades para a “casa” mais valiosa). Vamos ver o número  $1001110010_{(2)}$  (note que ele tem apenas 10 dígitos):

|                           |
|---------------------------|
| 1001110010 <sub>(2)</sub> |
|---------------------------|

|   |     |     |     |
|---|-----|-----|-----|
| 1 | 001 | 110 | 010 |
|---|-----|-----|-----|

|     |     |     |     |
|-----|-----|-----|-----|
| 001 | 001 | 110 | 010 |
|-----|-----|-----|-----|

|   |   |   |   |
|---|---|---|---|
| 1 | 1 | 6 | 2 |
|---|---|---|---|

# Resultado: 1162<sub>(8)</sub>

“Ei João, na verdade não é necessário adicionar os zeros à esquerda. Afinal, dá para perceber que 001 e 1 são a mesma coisa – convertidos para octal valem 1! Assim como 10 e 010 são a mesma coisa, pois convertidos para octal valem 2!”

Precisamente, caro leitor! Só mostrei desse jeito para que fosse mais fácil entender. Como no exemplo anterior, em que o 1 ficou sobrando no início do número, era fácil deduzir que o 1<sub>(2)</sub> é a mesma coisa que 1<sub>(8)</sub> (e, claro, é o mesmo que 1<sub>(10)</sub> e 1<sub>(16)</sub> também!).

Portanto, na tabela anterior, o número 1001110010 na base binária é a mesma coisa que 1162 na base octal.

Converter da base octal para a binária requer a mesma tabela usada nos dois exemplos anteriores, apenas com a inversão das posições dos números inicial e final. Vamos, como exemplo, converter o número 6741 da base octal para seu equivalente em binário.

## 6741<sub>(8)</sub>

|     |     |     |     |
|-----|-----|-----|-----|
| 6   | 7   | 4   | 1   |
| 110 | 111 | 100 | 001 |

## Resultado: 110111100001<sub>(2)</sub>

Segundo o que foi visto, o número octal 6741 é equivalente ao número binário 110111100001.

Veja agora um exemplo de um número que vai apresentar, em seu início, um grupo incompleto de três dígitos binários e, claro, só para lembrar, não será necessário adicionar zeros, afinal, zero à esquerda não vale nada! Analisemos o número 1332 na base octal.

## 1332<sub>(8)</sub>

|   |   |   |   |
|---|---|---|---|
| 1 | 3 | 3 | 2 |
|   |   |   |   |

1" 011" 011" 010"  
**Resultado:** 1011011010(2)

Portanto, o número octal 1332 é equivalente ao número binário 1011011010.

#### 12.3.4. Da base binária para a hexadecimal (e vice-versa)

De forma análoga à base octal, podemos fazer uma relação direta entre os dígitos hexadecimais e grupos de dígitos binários. Isso é possível nas bases octal e hexadecimal porque 8 e 16 são potências de 2 (base do sistema binário). Mas, sem filosofar, lá vai a tabela:

| Binário | Hexa | Binário |  |
|---------|------|---------|--|
| 0000    | 0    | 1000    |  |
| 0001    | 1    | 1001    |  |
| 0010    | 2    | 1010    |  |
| 0011    | 3    | 1011    |  |
| 0100    | 4    | 1100    |  |
| 0101    | 5    | 1101    |  |
| 0110    | 6    | 1110    |  |
| 0111    | 7    | 1111    |  |

Note que, cada dígito hexadecimal é representado como um grupo de quatro dígitos binários (bits). Portanto, para proceder com uma conversão de base 2 para base 16, é necessário,

primeiramente, separar o número binário em grupos de quatro dígitos. Lembre-se: comece a separação pela direita (unidades) para, quando chegar ao início do número, verificar a necessidade ou não, de adicionar zeros extras.

Analisemos o número  $100111010001_{(2)}$ . A conversão para hexadecimal se dará assim:

|                         |      |      |
|-------------------------|------|------|
| $100111010001_{(2)}$    |      |      |
| 1001                    | 1101 | 0001 |
| 9                       | D    | 1    |
| Resultado: $9D1_{(16)}$ |      |      |

Mais uma vez é necessário memorizar a tabela que associa binários a hexadecimais? Se for possível, sim! Mas, se você é avesso a qualquer tipo de “decoreba”, é só usar (novamente) o processo de conversão de binário para decimal em cada grupo de 4 dígitos, lembrando, claro, que, depois de convertido para decimal, se o valor for 10 vale A, 11 é B, 12 é C, 13 é D, 14 é E e, finalmente, 15 é F.

Um exemplo, no segundo grupo mostrado no exemplo anterior (1101), podemos ver que as casas  $2^3$ ,  $2^2$  e  $2^0$  estão válidas ( $2^1$  tem 0) e que isso vale  $8+4+1$ , portanto,  $13_{(10)}$ . Ora, o número  $13_{(10)}$  vale  $D_{(16)}$ . Esclarecido?

Outra conversão: o número  $101111101101011$  da base 2 será convertido para a base 16:

|                          |      |      |     |
|--------------------------|------|------|-----|
| $101111101101011_{(2)}$  |      |      |     |
| 101                      | 1111 | 0110 | 101 |
| 5                        | F    | 6    | B   |
| Resultado: $5F6B_{(16)}$ |      |      |     |

Finalmente, o número 10111101101011 da base 2 é equivalente ao número 5F6B na base hexadecimal.

Vamos continuar convertendo agora o número 12F<sub>(16)</sub> para a base binária. Os zeros à esquerda do grupo mais à esquerda (representando o primeiro dígito hexadecimal) não precisam ser escritos, como você já sabe:

|                                     |      |      |
|-------------------------------------|------|------|
| 12F <sub>(16)</sub>                 |      |      |
| 1                                   | 2    | F    |
| 1                                   | 0010 | 1111 |
| Resultado: 100101111 <sub>(2)</sub> |      |      |

Como se pôde verificar, o número 12F na base 16 é igual a 100101111 na base 2.

Como mais um simples exemplo, segue a análise do número BA26 e sua conversão para a base 2:

|                                            |      |      |    |
|--------------------------------------------|------|------|----|
| BA26 <sub>(16)</sub>                       |      |      |    |
| B                                          | A    | 2    | 6  |
| 1011                                       | 1010 | 0010 | 01 |
| Resultado: 1011101000100110 <sub>(2)</sub> |      |      |    |

Logo, conclui-se que o número BA26<sub>(16)</sub> é equivalente a 1011101000100110<sub>(2)</sub>.

#### 12.3.5. Da base octal para a hexadecimal (e vice-versa)

Eis uma coisa muito incomum! Se não usamos, de jeito nenhum, essas duas bases, por que

fariamos uma conversão entre elas? Pois é! É como fazer câmbio entre dólares canadenses e ienes (moeda japonesa) – não existe nenhum interesse em fazer tal operação (para nosso cotidiano).

Não há uma forma muito fácil de converter diretamente números na base 8 para a base 16 e vice-versa, mas há um jeito fácil se houver uma “conversão intermediária”. A melhor maneira que posso sugerir é converter o número (seja ele octal ou hexa) para binário pelo método correto (já apresentado) e, depois disso, converter o número binário para o resultado final (hexa ou octal).

Exemplos? Converteremos o número octal 1733 para hexadecimal:

|                                   |      |      |     |
|-----------------------------------|------|------|-----|
| 1733 <sub>(8)</sub>               |      |      |     |
| 1                                 | 7    | 3    | 3   |
| 001                               | 111  | 011  | 011 |
| Meio: 001111011011 <sub>(2)</sub> |      |      |     |
| 0011                              | 1101 | 1011 |     |
| 3                                 | D    | B    |     |
| Resultado: 3DB <sub>(16)</sub>    |      |      |     |

Note que houve uma conversão inicial de octal para binário (separando os dígitos octais em grupos de três dígitos binários, como havíamos visto) e, por fim, uma conversão de binário da hexadecimal (criando um grupo de quatro dígitos binários para cada dígito hexa). O número 1733<sub>(8)</sub> é equivalente ao número 3DB<sub>(16)</sub>.

Outro exemplo de conversão octal-hexadecimal: vamos tentar converter o número 25612<sub>(8)</sub>:

|                      |  |  |  |
|----------------------|--|--|--|
| 25612 <sub>(8)</sub> |  |  |  |
|----------------------|--|--|--|



2

5

6

1

010

101

110

0

Meio: 010101110001010<sub>(2)</sub>

010

1011

1000

2

B

8

Resultado: 2B8A<sub>(16)</sub>

Pelo que se viu, o número 25612<sub>(8)</sub> é equivalente ao número 2B8A<sub>(16)</sub>.

Vamos agora converter um número hexadecimal para octal através do mesmo processo: convertê-lo primeiro para binário para, depois, converter o binário finalmente para octal. E o número sorteado é... BB19<sub>(16)</sub>. (Note que é BB19, e não BBB19 – até porque eu não desejo que haja um BBB19!!!)

BB19<sub>(16)</sub>

B

B

1

1011

1011

0001

Meio: 1011101100011001<sub>(2)</sub>

1

011

101

100

|                                  |   |   |   |
|----------------------------------|---|---|---|
| 1                                | 3 | 5 | 4 |
| Resultado: 135431 <sub>(8)</sub> |   |   |   |

O número BB19 da base 16 equivale ao número 135431 da base 8.

Vamos ao segundo e derradeiro exemplo de conversão hexadecimal-octal. Dessa vez vamos converter o número 9812<sub>(16)</sub>. Note que esse número não tem letras; portanto, a única coisa que o caracteriza como hexadecimal é o índice entre parênteses.

|                      |
|----------------------|
| 9812 <sub>(16)</sub> |
|----------------------|

|   |   |   |
|---|---|---|
| 9 | 8 | 1 |
|---|---|---|

|      |      |      |
|------|------|------|
| 1001 | 1000 | 0001 |
|------|------|------|

|                                       |
|---------------------------------------|
| Meio: 1001100000010010 <sub>(2)</sub> |
|---------------------------------------|

|   |     |     |     |  |
|---|-----|-----|-----|--|
| 1 | 001 | 100 | 000 |  |
|---|-----|-----|-----|--|

|   |   |   |   |  |
|---|---|---|---|--|
| 1 | 1 | 4 | 0 |  |
|---|---|---|---|--|

|                                  |
|----------------------------------|
| Resultado: 114022 <sub>(8)</sub> |
|----------------------------------|

E, para finalizar toda essa confusão, o número 9812 da base hexadecimal pode ser escrito como 114022 na base octal.

#### 12.3.6. Atenção ao índice do número

“Por que alguns números vêm com índice e outros não apresentam essa informação?”

Caro leitor, o índice é apenas um indicativo da base numérica na qual aquele número foi escrito. É uma informação necessária na maioria dos casos. Ela serve para evitar a ambiguidade

da interpretação de um número.

Veja, por exemplo, o número 101001. Ele está em que base?

“Em binário, claro!”

Não necessariamente, caro leitor! Os caracteres 1 e 0 são usados em todas as bases numéricas! Esse poderia ser o número 101.001 da nossa base decimal (cento e um mil e um). Esse número também pode ser octal ou hexadecimal.

É claro que o número  $101001_{(2)}$  não tem o mesmo valor que  $101001_{(8)}$ , nem  $101001_{(10)}$ , nem  $101001_{(16)}$ . São todos números bem diferentes!

Portanto, na maioria dos casos, não é suficiente escrever o número. É necessário também indicar a sua base! Esta é a razão do índice: evitar interpretações erradas!

Em apenas um único caso não se faz necessária a apresentação do índice: nos números hexadecimais que apresentam letras! Afinal, as letras A, B, C, D, E e F só aparecem na base hexadecimal! Portanto, ao ler o número 3AD9, dá para saber imediatamente que se trata de um número na base 16.

Mas mesmo em números hexadecimais é necessário usar, algumas vezes, o índice: quando o número não tem as letras (como no último exemplo de conversão). Ao escrever o número 9162, por exemplo, é necessário dizer se ele é hexadecimal ( $9162_{(16)}$ ) para não confundi-lo com o número 9162 decimal (nove mil cento e sessenta e dois da nossa base) – que seria escrito  $9162_{(10)}$ .

### 12.3.7. Operações aritméticas em bases diferentes

Apesar de não ser uma coisa muito comum, exceto para o pessoal de TI, algumas vezes, as bancas colocam questões que pedem operações aritméticas com números de bases variadas.

“E aí, João, se aparecer? Como fazer?”

O jeito mais simples de resolver isso, amigo leitor, é o mais trabalhoso. Sugiro que você converta os números solicitados para a base 10 (decimal), realize a operação aritmética do jeito que aprendeu há muito tempo (soma, subtração, multiplicação ou divisão) para achar um resultado também decimal (base 10). Depois, é só converter o resultado para a base que está sendo pedida pela questão.

Qual o resultado da soma de  $1455_{(8)}$  e  $1245_{(8)}$ ? Apresente o resultado também na base octal.

- **Primeiro passo:** converta-os para a base decimal (é mais fácil entender assim)! Não vou mostrar aquela tabela grande de novo, só vou informar que  $1455_{(8)}$  é igual a  $813_{(10)}$  e  $1245_{(8)}$  é igual a  $677_{(10)}$ .
- **Segundo passo:** some-os! O resultado de  $813+677$  dá 1490 (na base decimal).
- **Terceiro e último passo:** converta o resultado para a base desejada (no nosso caso, devemos converter o resultado para a base octal). A resposta final é  $2722_{(8)}$ .

Vamos a mais um teste? Que tal resolver  $100110_{(2)} \times 722_{(8)}$  e dar o resultado em hexadecimal?

- **Primeiro passo:** converta os dois operandos da questão para a base decimal.

$100110$  em binário é o equivalente a 38 na base decimal.

722 em octal é o mesmo que 466 em decimal.

- **Segundo passo:** realize a operação que a questão pede com os números já convertidos para decimal ( $38 \times 466$ ). O resultado é 17708. Lembre-se: de que isso está na base decimal!
- **Terceiro passo:** converta o resultado (que foi descoberto em decimal) para a base que a questão pede (no caso, hexadecimal). O resultado da questão, em hexadecimal, como ela pediu, é 452C.

Vamos seguir com um estudo que é muito comum em provas para o cargo de Auditor-Fiscal em diversas escalas: operações com a base 2 (operações lógicas).

#### 12.4. Operações lógicas na base 2 (noções de álgebra booleana)

Algumas das vantagens da base binária em relação às demais (incluindo a nossa famosa decimal) estão relacionadas às operações que podem ser realizadas com esses números. Há algumas operações que são realizadas somente na base 2, como as operações conhecidas como lógicas: AND (E), OR (OU) e NOT (NÃO) além de outras menos exigidas.

Muitos dos conceitos que vamos ver aqui podem ser aplicados, inclusive, a questões de raciocínio lógico (assunto também visto em concursos públicos em geral). Depois darei mais exemplos acerca disso.

Lembre-se: como vamos lidar com operações lógicas, os números envolvidos são apenas 0 e 1 (base binária), que podemos simplesmente chamar de bits. (Lembra-se, no início deste livro? Bit significa dígito binário).

##### 12.4.1. Operador AND (E)

Uma operação AND entre dois bits resulta em 1 apenas se os dois operandos forem 1. Ou seja, a operação AND apresenta a seguinte tabela-verdade com dois bits (operação  $x$  AND  $y$ ):

| x | y | $x$ AND $y$ |
|---|---|-------------|
| 0 | 0 | 0           |
| 0 | 1 | 0           |
| 1 | 0 | 0           |
| 1 | 1 | 1           |

O que podemos concluir? Se for usada a operação AND em dois bits quaisquer, o resultado da

operação será sempre 0, exceto quando os dois bits de origem forem 1.

“Ei, João, e se a operação AND for realizada entre dois números binários com mais de um bit? Já vi isso em prova uma vez!”

Deve-se analisar bit a bit. Coloque-os organizadamente em uma tabela (ambos os números) com os bits organizados em casas (comece preenchendo da direita para a esquerda, porque se um dos números tiver menos bits, a gente completa com zero à esquerda). O exemplo é a operação 01110100 AND 11100010, que pode ser resolvida assim:

|                              |                 |   |   |  |
|------------------------------|-----------------|---|---|--|
| Operação:                    | 01110100 AND 11 |   |   |  |
| Bits do<br>Número<br>1:      | 0               | 1 | 1 |  |
| Bits do<br>Número<br>2:      | 1               | 1 | 1 |  |
| Bits<br>depois do<br>AND:    | 0               | 1 | 1 |  |
| Resultado<br>da<br>Operação: | 01100000        |   |   |  |

Note que o primeiro número apresenta um 0 à esquerda, que, tecnicamente, não é necessário ser escrito.

Vamos testar mais uma vez com a operação 11100101 AND 10110101:

|                        |                       |   |   |  |
|------------------------|-----------------------|---|---|--|
| Operação:              | 11100101 AND 10110101 |   |   |  |
| Bits do Número 1:      | 1                     | 1 | 1 |  |
| Bits do Número 2:      | 1                     | 0 | 1 |  |
| Bits depois do AND:    | 1                     | 0 | 1 |  |
| Resultado da Operação: | 10100101              |   |   |  |

#### 12.4.2. Operador OR (OU)

Uma operação OR entre dois bits resulta em 1 se pelo menos um dos bits for 1. A tabela-verdade da operação  $x \text{ OR } y$  (sendo  $x$  e  $y$  bits quaisquer) é esta:

| x | y | x OR y |
|---|---|--------|
| 0 | 0 | 0      |
| 0 | 1 | 1      |
| 1 | 0 | 1      |
| 1 | 1 | 1      |

E então? O resultado de uma operação OR será sempre 1, exceto se os dois bits forem 0.

De forma análoga à operação AND com números formados por vários bits, na operação OR devemos analisar um bit por vez para obter o resultado da operação: analisemos a operação 01100110 OR 10001100:

| Operação:               | 01100110 OR 10001100 |   |   |  |
|-------------------------|----------------------|---|---|--|
| Bits do<br>Número<br>1: | 0                    | 1 | 1 |  |
| Bits do<br>Número<br>2: | 1                    | 0 | 0 |  |
| Bits                    |                      |   |   |  |

|           |          |   |   |
|-----------|----------|---|---|
| depois do | 1        | 1 | 1 |
| OR        |          |   |   |
| Resultado | 11101110 |   |   |
| da        |          |   |   |
| Operação: |          |   |   |

Continuando com o segundo exemplo: 10001010 OR 00101001.

|                          |                      |   |   |  |
|--------------------------|----------------------|---|---|--|
| Operação:                | 10001010 OR 00101001 |   |   |  |
| Bits do<br>Número<br>1:  | 1                    | 0 | 0 |  |
| Bits do<br>Número<br>2:  | 0                    | 0 | 1 |  |
| Bits<br>depois do<br>OR: | 1                    | 0 | 1 |  |
| Resultado                |                      |   |   |  |



da

10101011

Operação:

**12.4.3. Operador NOT (Não)**

O operador de NEGAÇÃO (NOT) é unário, ou seja, não está no meio de uma operação entre dois operandos; esse operador precede um único operando, invertendo seu bit. A tabela-verdade desse operador é mais simples:

| x | NOT x |
|---|-------|
| 0 | 1     |
| 1 | 0     |

Fazer a operação NOT com um número grande também requer que façamos a operação com um bit por vez, como na operação NOT 01100111.

| Operação:           | NOT 01100111 |   |   |  |
|---------------------|--------------|---|---|--|
| Bits do Número:     | 0            | 1 | 1 |  |
| Bits depois do NOT: | 1            | 0 | 0 |  |
| Resultado           |              |   |   |  |

da

10011000

Operação:

**12.4.4. Operador XOR (ou exclusivo)**

Esse é “um dos outros” operadores lógicos que podemos usar em números binários. Estava aqui “matutando” se o colocaria ou não, e acabei vencido pela minha consciência, que disse: “Rapaz, e se isso cai na prova, como os alunos vão fazer?”

Essa operação só pode resultar em 1 se os valores dos bits forem diferentes e resulta em 0 quando os dois bits operados são iguais. A tabela para x XOR y é esta:

| X | y | x XOR y |
|---|---|---------|
| 0 | 0 | 0       |
| 0 | 1 | 1       |
| 1 | 0 | 1       |
| 1 | 1 | 0       |

Vamos analisar a operação de OU Exclusivo com números binários grandes. Como nosso exemplo, tomemos 01001101 XOR 10011100.

| Operação:         | 01001101 XOR 10011100 |   |   |  |
|-------------------|-----------------------|---|---|--|
| Bits do Número 1: | 0                     | 1 | 0 |  |

|                              |          |   |   |  |
|------------------------------|----------|---|---|--|
| Bits do<br>Número<br>2:      | 1        | 0 | 0 |  |
| Bits<br>depois do<br>XOR:    | 1        | 1 | 0 |  |
| Resultado<br>da<br>Operação: | 11010001 |   |   |  |

#### 12.4.5. Algumas regras gerais

Não gosto muito de regras para memorizar, mas aqui vão algumas que podem ser comprovadas com tudo o que vimos. Essas regrinhas seguirão como resumos de todo o assunto visto. Tomemos, então, A como um bit (ou seja, pode assumir 0 ou 1):

|                            |                           |
|----------------------------|---------------------------|
| $A \text{ AND } 0 = 0$     | $A \text{ AND } 1 = A$    |
| $A \text{ OR } 0 = A$      | $A \text{ OR } 1 = 1$     |
| $A \text{ AND NOT}(A) = 0$ | $A \text{ OR NOT}(A) = 1$ |
| $A \text{ AND } A =$       |                           |

$A$  $A \text{ OR } A = A$  $A \text{ XOR } A = 0$  $A \text{ XOR } \text{NOT}(A) = 1$ 

Essas regrinhas podem ser usadas para responder facilmente a algumas questões de raciocínio lógico e outras de matemática binária mesmo. Tente verificar o porquê da veracidade dessas operações.

#### 12.4.6. Aplicação prática da álgebra booleana (endereço IP)

Finalmente, chegamos a um ponto em que conseguimos entender vários conceitos necessários para a resolução de um problema muitas vezes proposto em provas que envolvem conhecimentos avançados de endereçamento IP.

(Se o seu estudo não precisa de Redes de Computadores e/ou conhecimentos muito avançados sobre endereço IP, nem leia este tópico! Mas se você precisa desses assuntos, é bom que já os tenha lido para poder entender o que está sendo tratado aqui.)

Esse nível de questões, normalmente, é encontrado em provas de TI (provas para o pessoal de Informática), mas nada impede que algum dia, se estiver de mau humor, a ESAF ouse colocar esse tipo de questão em uma prova de Auditor-Fiscal. (Como se um candidato a auditor já não tivesse muita coisa para estudar.)

A questão é: sabe como identificar o ID da rede em um endereço IP? Faça uma operação de AND com a máscara de sub-rede usada por aquele micro!

Assim: imaginemos um micro 192.168.14.235/20 (micro 1) e outro 192.168.51.99/20 (micro 2). Se fosse pedido para você determinar se eles fazem parte da mesma rede, o que você faria?

“Pularia a questão, João, e partiria para uma menos complicada.”

É, ótima solução, mas se você realmente estiver interessado, caro leitor, em resolver a questão, aqui seguem os passos:

Primeiro passo: converta os endereços IP e a máscara para binário.

Micro 1: 192.168.14.235 – 11000000.10101000.00001110.11101011

Máscara: /20 (20 bits 1) – 11111111.11111111.11110000.00000000 – seria o equivalente a 255.255.240.0 (mas essa forma de escrita não é necessária).

Segundo passo: descobrir o ID da rede deste micro. Para isso, fazemos um AND entre o IP e a máscara. Assim:

(IP do

|                                 |       |
|---------------------------------|-------|
| 1 1 0 0 0 0 0 0 1 0 1 0 1 0 0 0 | micro |
| 0 0 0 0 1 1 1 0 1 1 1 0 1 0 1 1 | 1)    |

|                                 |              |
|---------------------------------|--------------|
| 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 | (Máscara)    |
| 1 1 1 1 1 0 0 0 0 0 0 0 0 0 0 0 |              |
| 0 0                             |              |
| 1 1 0 0 0 0 0 0 1 0 1 0 1 0 0   | (ID da rede) |
| 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0   |              |
| 0 0                             |              |

Separando o ID da rede em quatro grupos de 8 bits (octetos), temos:

11000000.10101000.00000000.00000000 (ID da rede)

Convertendo cada octeto separadamente para a base decimal, conclui-se que a rede à qual pertence o micro 192.168.14.235/20 chama-se 192.168.0.0/20.

Terceiro passo: descobrir o endereço IP do micro 2 (192.168.51.99/20) em binário (a máscara é a mesma do micro 1, portanto será só repetir).

Micro 2: 192.168.51.99 – 11000000.10101000.00110011.01100011

Máscara: /20 (20 bits 1) – 11111111.11111111.11110000.00000000

Quarto passo: realizar um AND entre IP do micro 2 e máscara.

|                               |                 |
|-------------------------------|-----------------|
| 1 1 0 0 0 0 0 0 1 0 1 0 1 0 0 | (IP do micro 2) |
| 0 0 0 1 1 0 0 1 1 0 1 1 0 0 0 |                 |
| 1 1                           |                 |

|                                                                       |              |
|-----------------------------------------------------------------------|--------------|
| 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1<br>1 1 1 1 1 0 0 0 0 0 0 0 0 0 0<br>0 0 | (Máscara)    |
| 1 1 0 0 0 0 0 0 1 0 1 0 1 0 0<br>0 0 0 1 1 0 0 0 0 0 0 0 0 0 0<br>0 0 | (ID da rede) |

Separando o ID da rede em quatro grupos de 8 bits (octetos), temos:

11000000.10101000.00110000.00000000 (ID da rede)

Finalmente, convertendo cada octeto separadamente para a base decimal, dá para perceber que a rede do micro 2 chama-se 192.168.48.0/20.

Portanto, os dois micros apresentados não pertencem à mesma rede!

## 12.5. Considerações finais

Espero que este assunto, muitas vezes não atraente e não relevante, seja entendido da melhor maneira, caro leitor e querido aluno. Sei que não é comum encontrar questões com esses assuntos, mas, se aparecerem em alguma prova, você já estará preparado!

Fique sempre com Deus!

João Antonio

## 12.6. Questões de Aritmética Computacional

1. O número  $123_{(10)}$  será escrito, na base binária, como:
  - a) 0110110;
  - b) 1111011;
  - c) 1001010;
  - d) 1110000;
  - e) 1010001.
2. O número  $255_{(10)}$  será escrito, na base hexadecimal, como:
  - a) FA;
  - b) FF;
  - c) CF;
  - d) 22;
  - e) 9F.
3. Qual o maior número decimal que pode ser representado com um número binário de 8 bits?
  - a) 256;
  - b) 512;
  - c) 1024;
  - d) 255;
  - e) 64.
4. O número  $212_{(10)}$  será escrito, na base octal, como:
  - a) 192;
  - b) 323;
  - c) 328;
  - d) 122;
  - e) 324.
5. O número  $F5AB_{(16)}$  será escrito, na base binária, como:
  - a) 1111010110101011;
  - b) 1001111101010001;
  - c) 1010101010101010;
  - d) 1010001110110000;
  - e) 1001111001100010.
6. A operação  $34B - 23A$  resulta no número hexadecimal:
  - a) 121;
  - b) 351;
  - c) 4AB;
  - d) 111;

e) 1A1.

7. A operação  $123_{(8)} + 103_{(8)}$  resulta no número decimal:

- a) 226;
- b) 150;
- c) 151;
- d) 3AB;
- e) 101.

8. Julgue os itens a seguir como Certos ou Errados de acordo com a conversão realizada.

- a)  $123_{(16)} = 231_{(10)}$
- b)  $234_{(16)} = 564_{(10)}$
- c)  $3AB_{(16)} = 939_{(10)}$
- d)  $BC2_{(16)} = 3011_{(10)}$
- e)  $1101_{(8)} = 1001000001_{(2)}$
- f)  $1277_{(8)} = 10010111111_{(2)}$
- g)  $621_{(8)} = 110010001_{(2)}$
- h)  $345_{(8)} = 11100101_{(2)}$
- i)  $101111110_{(2)} = 376_{(8)}$
- j)  $111100011_{(2)} = 221_{(8)}$
- k)  $100000001_{(2)} = 101_{(8)}$
- l)  $10001111_{(2)} = 217_{(8)}$
- m)  $237_{(10)} = 01100110_{(2)}$
- n)  $345_{(10)} = 101011001_{(2)}$
- o)  $344_{(10)} = 101011000_{(2)}$
- p)  $86_{(10)} = 100100100_{(2)}$
- q)  $AB12_{(16)} = 23875_{(8)}$
- r)  $FE1C_{(16)} = 177034_{(8)}$
- s)  $1122_{(16)} = 10442_{(8)}$
- t)  $12AA_{(16)} = 11252_{(8)}$

9. Julgue os itens a seguir como Certos ou Errados de acordo com a operação aritmética.

- a)  $0101_{(8)} + 1212_{(8)} = 1313_{(8)}$
- b)  $4532_{(8)} + 2112_{(8)} = 6644_{(8)}$
- c)  $103_{(16)} + 2AB_{(16)} = 3AE_{(16)}$
- d)  $AABB_{(16)} + ACAD_{(16)} = FACA_{(16)}$
- e)  $01100110_{(2)} + 10011101_{(2)} = 10011111_{(2)}$



f)  $104_{(8)} + 345_{(16)} = 32480_{(8)}$

g)  $234_{(16)} + AA1_{(16)} = 6325_{(8)}$

h)  $110_{(8)} + FF_{(16)} = 101000111_{(2)}$

i)  $123_{(10)} + 657_{(8)} = 22A_{(16)}$

## Gabaritos

### Capítulo 2

#### Estilo FCC

|      |      |      |      |      |
|------|------|------|------|------|
| 1. D | 2. C | 3. A | 4. E | 5. A |
|------|------|------|------|------|

#### Estilo ESAF

|      |      |      |      |
|------|------|------|------|
| 1. C | 2. E | 3. A | 4. C |
|------|------|------|------|

### Capítulo 4

|      |      |      |      |      |      |      |
|------|------|------|------|------|------|------|
| 1. D | 2. B | 3. E | 4. A | 5. E | 6. D | 7. C |
|------|------|------|------|------|------|------|

### Capítulo 6

|      |      |      |      |       |
|------|------|------|------|-------|
| 1. d | 2. b | 3. c | 4. a | 5. b  |
| 6. d | 7. d | 8. b | 9. a | 10. d |

### Capítulo 7

|      |      |      |      |      |      |      |
|------|------|------|------|------|------|------|
| 1. E | 2. C | 3. D | 4. B | 5. D | 6. C | 7. E |
|------|------|------|------|------|------|------|

## Capítulo 8

|      |      |      |      |      |      |
|------|------|------|------|------|------|
| 1. D | 2. E | 3. C | 4. C | 5. C | 6. D |
|------|------|------|------|------|------|

## Capítulo 9

|      |      |      |      |      |      |      |
|------|------|------|------|------|------|------|
| 1. D | 2. D | 3. E | 4. A | 5. C | 6. B | 7. A |
|------|------|------|------|------|------|------|

## Capítulo 10

|      |      |      |      |      |      |      |
|------|------|------|------|------|------|------|
| 1. D | 2. C | 3. C | 4. D | 5. D | 6. C | 7. A |
|------|------|------|------|------|------|------|

## Capítulo 11

|      |      |      |      |      |
|------|------|------|------|------|
| 1. B | 2. C | 3. B | 4. E | 5. A |
| 6. D | 7. C | 8. A | 9. E |      |

## Capítulo 12

|      |      |      |      |      |      |      |
|------|------|------|------|------|------|------|
| 1. B | 2. B | 3. D | 4. E | 5. A | 6. D | 7. A |
|------|------|------|------|------|------|------|

|    |  |
|----|--|
| 8. |  |
|----|--|

|                         |                          |
|-------------------------|--------------------------|
| a) Errado (291)         | b) Certo                 |
| c) Certo                | d) Errado (3010)         |
| e) Certo                | f) Errado<br>(101011111) |
| g) Certo                | h) Certo                 |
| i) Errado (576)         | j) Errado (743)          |
| k) Errado (401)         | l) Certo                 |
| m) Errado<br>(11101101) | n) Certo                 |
| o) Certo                | p) Errado<br>(1010110)   |
| q) Errado<br>(125422)   | r) Certo                 |
|                         |                          |

|          |          |
|----------|----------|
| s) Certo | t) Certo |
|----------|----------|

|                       |                                  |
|-----------------------|----------------------------------|
| 9)                    |                                  |
| a) Certo              | b) Certo                         |
| c) Certo              | d) Errado (15768 em hexadecimal) |
| e) Errado (100000011) | f) Errado (1611 em octal)        |
| g) Certo              | h) Certo                         |
| i) Certo              |                                  |

## Bibliografia

ALBUQUERQUE, Fernando. *TCP/IP – Internet: protocolos e tecnologia*. 3. ed. Rio de Janeiro: Axcell Books, 2001.

COMER, Douglas E. *Rede de computadores e internet*. Trad. Marinho Barcellos. 2. ed. Porto Alegre: Bookman, 2001.

PINHO, Roberto N. Lima Caribe. *Introdução à computação*. São Paulo: FTD, 1996.

SOUZA, Lindeberg Barros de. *Redes Cisco CCNA – faça certificação*. São Paulo: Érica, 2002.

TANEMBAUM, Andrew S. *Redes de computadores*. Trad. da 3. edição. Insight Serviços de Informática. Rio de Janeiro: Campus, 1997.

\_\_\_\_\_. *Organização estruturada de computadores*. Rio de Janeiro: LTC, 2001.

THING, Lowell. *Dicionário de tecnologia*. Trad. Bazán Tecnologia e Linguística e Texto Digital. São Paulo: Futura, 2003.

VASCONCELOS, Laércio. *Hardware total*. São Paulo: Makron Books, 2002.

## SITES

[www.guiadohardware.net](http://www.guiadohardware.net) (autor: Carlos E. Morimoto)

[www.clubedohardware.com.br](http://www.clubedohardware.com.br) (autor: Gabriel Torres)

[www.julioibatisti.com.br](http://www.julioibatisti.com.br) (autor: Julio Batisti)

[www.laercio.com.br](http://www.laercio.com.br) (autor: Laércio Vasconcelos)

[www.tomshardware.com](http://www.tomshardware.com)

[www.intel.com.br](http://www.intel.com.br)

[www.amd.com.br](http://www.amd.com.br)

[www.cisco.com](http://www.cisco.com)

[www.microsoft.com.br](http://www.microsoft.com.br)

[www.baboo.com.br](http://www.baboo.com.br)

[www.wikipedia.org](http://www.wikipedia.org) (Qual o problema? As bancas também usam...)

## MARCAS REGISTRADAS

Windows, Word, Excel, Access, PowerPoint, SQL Server são marcas registradas da Microsoft Corporation.

Pentium, Celeron, Centrino, HT, Itanium são marcas registradas da Intel Corporation.

Athlon, Duron, Athlon 64 são marcas registradas da AMD Corporation.

Oracle é marca registrada da Oracle Corporation.

Demais produtos de hardware e software citados neste livro são marcas registradas de suas respectivas produtoras/fabricantes.